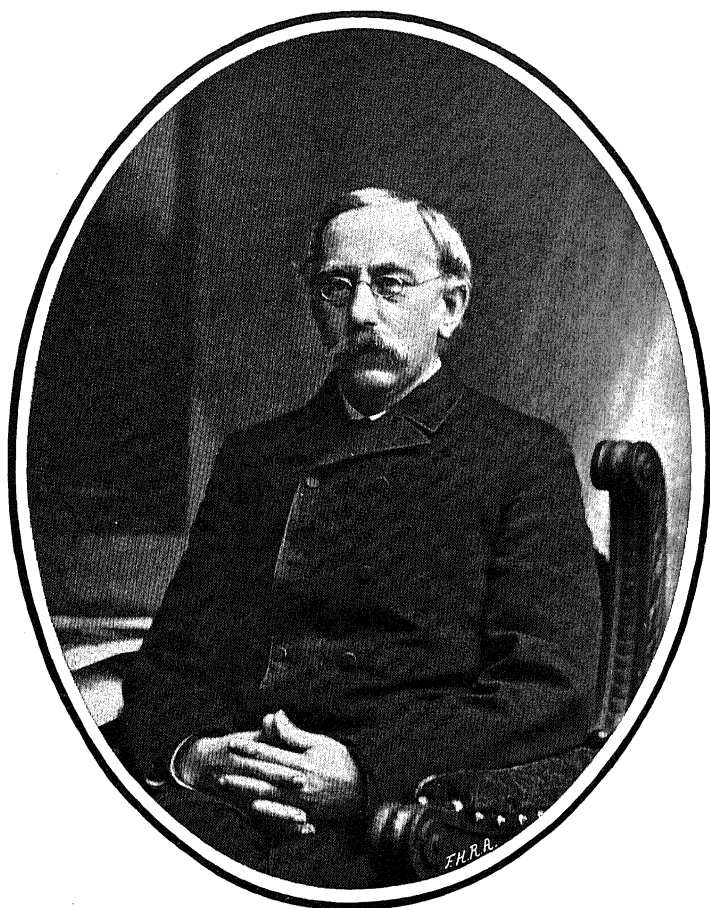


**This book is with
tight
Binding**

**CARNEGIE INSTITUTE
OF TECHNOLOGY**



THE LIBRARY



Loring

ŒUVRES SCIENTIFIQUES

DE

L. LORENZ

REVUES ET ANNOTÉES

PAR

H. VALENTINER

PUBLIÉES AUX FRAIS DE LA FONDATION CARLSBERG

TOME SECOND

DEUXIÈME FASCICULE

COPENHAGUE

LIBRAIRIE LEHMANN & STAGE

IMPRIMERIE DE BIANCO LUNO

1904

TABLE DES MATIÈRES DU TOME SECOND.

Mémoire sur la théorie de l'élasticité des corps homogènes à élasticité constante, p. 1.

Sur le nombre des molécules contenues dans un milligramme d'eau, p. 47.

Détermination du degré de chaleur en mesures absolues, p. 55.

La résistance électrique du mercure en mesures absolues, p. 85.

Sur les méthodes à employer pour la détermination de l'ohm, p. 117.

Détermination de la résistance électrique du mercure en mesures électromagnétiques absolues, p. 129.

Sur la propagation de l'électricité, p. 183.

Sur les conductibilités électrique et calorifique des métaux, p. 241.

Sur le développement des fonctions au moyen d'intégrales définies, p. 317.

Un théorème sur la fonction potentielle, p. 331.

Sur l'évaluation des aires, p. 357.

Sur le mouvement permanent d'un liquide, p. 367.

Sur la résolution des équations algébriques au moyen de séries et d'intégrales définies, p. 385.

Contribution à la la théorie des nombres, p. 401.

Sur la compensation des erreurs d'observation, p. 433.

Sur la réduction du facteur enlérien, p. 467.

Équations cinétiques fondamentales d'un système de points,
p. 483.

Sur le développement des fonctions arbitraires au moyen de
fonctions données, p. 493.

Sur les nombres premiers, p. 521.

Recherches analytiques sur les nombres de nombres premiers,
p. 531.

VIE ET TRAVAUX DE L.-V. LORENZ.

Ludvig Valentin Lorenz naquit le 18 janvier 1829, à Helsingør. Son père était de Stralsund et s'était marié avec la fille d'un boulanger de Helsingør, à qui il succéda dans son commerce. En 1835, Lorenz le père s'établit comme marchand à Maribo, petite ville située dans l'île Laaland. C'est pour cela que L.-V. Lorenz a reçu son instruction à l'école supérieure de Nykjøbing, petite ville voisine de Maribo. De très bonne heure, il fit voir un penchant extraordinaire pour l'étude des mathématiques. Lorenz le père s'intéressait lui-même aux mathématiques et à la mécanique et avait assez largement initié son fils à ces branches de la science, et notre Lorenz a de son propre chef continué ces études. Ainsi, et son éducation et ses dispositions personnelles ont fait de lui un autodidacte: il dit lui-même qu'il a pris l'habitude mauvaise de trouver, sans s'inquiéter des recherches des autres, la solution des problèmes, quitte à savoir ensuite qu'ils étaient déjà résolus depuis longtemps. Cette manière de faire ses études est assez lente, et tandis que d'un côté il avait de bonne heure développé ses dons pour les mathématiques, il était, d'un autre côté, en retard dans les autres branches de son instruction. On comprend qu'il ne pouvait pas profiter beaucoup de l'instruction mathématique donnée à l'école: il a dit plus tard qu'il n'avait pas en général tiré grand

profit de cette instruction; qu'il avait au contraire dans les années de l'école le plein loisir et tout le temps de lire des choses qui ne faisaient pas partie de l'enseignement. Après avoir passé son examen comme étudiant en 1849, il alla à l'école polytechnique de Copenhague pour préparer son examen de mécanicien; mais la manière dont Ramus (le professeur d'alors) exposait les mathématiques supérieures le rebuta. C'est pourquoi il se décida à étudier la chimie; il fit la plupart de ses études sans l'aide des professeurs, n'assista qu'à peu de leçons, mais s'attacha aux problèmes de la physique, surtout aux spéculations sur la nature de la chaleur.

Par suite il n'obtint à son examen qu'une note assez médiocre. Après cet examen, il devint précepteur chez le comte Moltke Huidtfeldt, puis il gagna sa vie en enseignant dans plusieurs écoles de Copenhague. De septembre 1858 jusqu'à juillet 1859, Lorenz séjournait à Paris aux frais de l'état. Il y a suivi les cours de Regnault, des deux Becquerel père et fils, de Despretz, de Desains, de Bertrand et de Lamé. De 1866 à 1887, Lorenz fut professeur au séminaire de Blaagaard, et de 1866 à 1887 professeur de physique à l'école royale militaire supérieure (plus tard nommée l'école des officiers). Cette dernière profession fut de sa part l'objet d'une préférence marquée et il avait à un haut degré l'approbation de ces collègues et de ses élèves, quoique ses leçons exigeassent assez souvent un trop grand effort des auditeurs. En 1887, la direction de la fondation Carlsberg lui offrit une situation de savant libre, qu'il accepta volontiers. Dès lors il pouvait faire ce qu'il avait désiré depuis longtemps: consacrer tout son temps à ses recherches scientifiques. Malheureusement, il ne

III

lui restait que peu de temps à travailler. Le 9 juin 1891, une apoplexie du cœur, conséquence d'une maladie dont il avait déjà longtemps souffert, mettait fin à sa vie. Lorenz s'était marié le 12 août 1862 avec Agathe Fogtmann. En 1876, il fut nommé professeur agrégé et en 1887, à son départ de l'école des officiers, conseiller d'état.

Depuis 1866, il était membre de l'académie royale des sciences de Copenhague et en 1887 il reçut le titre honorifique de docteur en philosophie de l'université d'Upsala.

Nous avons dit ci-dessus que Lorenz s'occupa de bonne heure des problèmes mathématiques. Son intérêt pour la physique fut éveillé par quelques leçons faites en 1840 par C. Carlsen, plus tard ingénieur des travaux maritimes. Mais, du reste, il s'intéressa ensuite très vivement à divers autres sujets. Il avait pendant quelque temps l'intention de se faire théologien; il étudia l'hébreu, sur lequel il passa un examen, et dans ses années d'étudiant il était plein de Søren Kirkegaard (philosophe et théologien danois). La politique aussi l'intéressa beaucoup, et il a écrit en 1865 un pamphlet sous ce titre „Notre lutte intérieure“. Ce pamphlet conserve encore quelque intérêt, car Lorenz y prédit certaines conséquences de la constitution, qui n'ont pas manqué de se manifester. Mais enfin son intérêt pour la physique mathématique prévalut, et il développa successivement la rare association des aptitudes expérimentales et mathématiques qu'il possédait.

Les productions de Lorenz sont presque toutes mathématiques en même temps que physiques; tantôt prévaut le côté mathématique, tantôt le côté physique;

mais presque toujours que les mathématiques lui servent à développer les idées physiques. Son intérêt pour la mathématique pure n'est pas aussi développé que son intérêt pour la physique, et, quand il traite un problème purement mathématique, il se soucie plus des résultats que de l'exactitude des raisonnements. (Voir, par exemple, le dernier mémoire „Sur le nombre des nombres premiers“.) Je crois qu'on peut dire que son talent principal était sa disposition brillante pour l'invention des méthodes, combinée avec un sens très fin de la valeur des résultats approchés, même quand il ne peut donner aucune limite des erreurs commises. C'est ce double don qui lui permet, par exemple, de parvenir au but dans le mémoire „Sur la lumière réfléchie et réfractée par une sphère transparente“.

Son défaut est dans ses raisonnements purement spéculatifs, qui sont très peu clairs. Aussi une grande partie de ses mémoires sont très difficiles à comprendre, surtout, quand les expériences ne peuvent pas confirmer ou infirmer les résultats (voir par exemple le mémoire: „Sur la compensation des erreurs d'observation“). J'ai dit plus haut que son intérêt pour la physique était plus grand que son intérêt pour la mathématique; je dois pourtant ajouter que quelques-uns de ses mémoires purement mathématiques conserveront toujours leur importance, par exemple le beau mémoire sur les fonctions besséliennes.

Cela dit, je considérerai un peu plus en détail les objets différents des recherches scientifiques de Lorenz.

Je commencerai par ses travaux optiques. Ces travaux sont assez nombreux et en même temps leur

étendue est considérable: ils remplissent, comme on le voit, tout le premier volume des „Œuvres“.

Les travaux de Lorenz sur la théorie de la lumière embrassent presque toute cette théorie et formeraient avec quelques suppléments une base excellente pour le développement de l'Optique physique tout entière. On peut indiquer en peu de mots le but des cinq premiers mémoires en question: ce but est d'établir une formule mathématique qui résume les résultats de toutes les recherches faites sur la théorie de la lumière. Lorenz veut de plus que cette formule soit déduite non de la considération des causes physiques des phénomènes lumineux, mais seulement des résultats d'observation, et, comme je l'ai dit, qu'elle comprenne tous les résultats, à l'exception toutefois de ceux qui dépendent de forces inconnues (électriques, chimiques, etc.). Inversement, des phénomènes embrassés par la formule on ne pourra déduire que les équations fondamentales et non une théorie physique (voir le cinquième mémoire: Sur la théorie de la lumière, p. 145). C'est peut-être dans les forces inconnues qu'est cachée cette théorie (voir le sixième mémoire). Il est évident que, pour établir cette formule, on a besoin d'un fondement expérimental, admissible en toute généralité. Lorenz le trouve dans la théorie du mouvement ondulatoire de la lumière et dans les formules de Fresnel relatives à l'intensité des rayons réfractés et réfléchis par la surface de séparation de deux milieux homogènes et isotropes.

Le but des trois premiers mémoires est de démontrer la légitimité des hypothèses admises et de définir avec plus de précision la manière dont s'opère le mouvement ondulatoire.

Le premier et le troisième mémoire* cherchent à déterminer la direction des vibrations, le premier par des considérations théoriques et par des expériences, le troisième uniquement par des calculs. On possède, comme on sait, deux théories également admissibles de la direction des vibrations lumineuses, celle de Fresnel et celle de Neumann. Lorenz cherche dans les deux mémoires en question à trancher la question de savoir laquelle est préférable. Je crois pourtant que ni ses calculs ni ses expériences ne sont décisifs. Les calculs ne donnent qu'une certaine approximation et les expériences ne concordent pas complètement avec les calculs. Peut-être est-il impossible de trancher la question sans définir avec plus de précision en quoi consiste le mouvement ondulatoire; car, d'après la théorie électromagnétique, le mot „éther“ est devenu insignifiant, et il faut définir si c'est par exemple la force électrique ou la force magnétique qui prend part au mouvement ondulatoire. Mais, quelle que soit la direction des vibrations, je ne crois pas qu'elle ait aucune influence essentielle sur la forme des équations fondamentales de la théorie de la lumière.

Le second mémoire du fascicule (Sur la réflexion de la lumière à la surface de séparation de deux milieux transparents et isotropes) traite de quelques expériences faites par Jamin sur la réflexion à la surface des corps transparents, expériences qui mettent en évidence quelques écarts par rapport aux formules de Fresnel. Dans ce

* Détermination de la direction des vibrations de l'éther lumineux par la polarisation de la lumière diffractée. — Détermination de la direction des vibrations de l'éther par la réflexion et par la réfraction de la lumière.

mémoire important, l'auteur cherche à démontrer que ces écarts ne sont qu'apparents et peuvent être expliqués par la supposition que deux corps ne sont jamais séparés par une surface proprement dite, mais qu'il existe toujours de l'un à l'autre une couche de transition, et que par conséquent l'indice de réfraction ne change pas brusquement quand on passe d'un corps à l'autre, mais varie toujours par degrés insensibles.

On sait que Christoffel (*Fortschritte der Physik*, t. 16, voir note 6) a fait à cette théorie l'objection, d'ailleurs fondée, qu'une certaine quantité que Lorenz suppose toujours très petite doit nécessairement devenir parfois infiniment grande. Mais cette quantité n'entre dans les formules de Lorenz que comme élément d'une intégrale, et, comme j'ai cherché à le démontrer, elle est vraisemblablement infinie de telle manière que l'intégrale soit toujours très petite, en sorte que l'objection indiquée ne peut suffire à renverser l'hypothèse de Lorenz.

Après avoir développé la base de sa théorie de la lumière dans les trois premiers mémoires, Lorenz établit dans le quatrième mémoire (*Sur la théorie de la lumière*) les équations fondamentales. Ces équations expriment le mouvement de la lumière dans un milieu quelconque, homogène ou hétérogène, isotrope ou anisotrope. La possibilité du développement de ces équations repose sur l'hypothèse admise par Lorenz, que les lois indiquées s'étendent à tous les cas sans exception, si petites que soient les dimensions des corps, fussent-elles même insignifiantes en comparaison de la longueur d'une onde lumineuse.

Lorenz se sert ensuite de ses équations fondamentales pour mettre en évidence comment on en peut

VIII

déduire quelques singularités bien connues de la lumière, à savoir la double réfraction et la rotation du plan de polarisation, en faisant quelques hypothèses simples sur la constitution moléculaire des corps.

Le développement de ces propriétés est très intéressant, car on reconnaît par les considérations de Lorenz qu'un corps doit être doublement réfringent, lorsqu'il est constitué par des couches périodiques. Des expériences ont confirmé cette conclusion. Un grand intérêt s'attache également à la théorie de la rotation du plan de polarisation, car Lorenz prouve que cette rotation doit avoir lieu quand le corps est constitué par des couches entrecroisées formant par leurs intersections mutuelles des figures géométriques semblables aux polyèdres réguliers.

Le cinquième mémoire donne plus d'extension à la théorie de la lumière. Il fait voir entre autres choses que la théorie de Lorenz, qui est fondée sur l'hypothèse de Fresnel, est en concordance avec celle de Neumann, si l'on fait l'hypothèse que les quantités au moyen desquelles est représenté le mouvement ondulatoire n'ont pas d'existence physique, mais sont simplement des grandeurs fictives, introduites pour faciliter le calcul des phénomènes directement observés.

On peut se demander si Lorenz a en effet atteint son but, qui était d'établir des formules mathématiques susceptibles d'embrasser, indépendamment de toute hypothèse physique, l'ensemble des résultats d'observation. A cette question je crois qu'on doit répondre qu'il n'y a pas complètement réussi. Car, même en supposant les lois de Fresnel absolument exactes, les expériences font seulement voir qu'elles s'appliquent à des surfaces de dimensions finies, tandis que les calculs de Lorenz

exigent qu'elles soient applicables à des surfaces infiniment petites, ce qui implique une hypothèse physique. Lorenz dit bien qu'il regarde toute autre formule, qui présente la même concordance avec les expériences, comme ayant la même validité (voir p. 91). Mais, en considérant la forme des équations, on reconnaîtra, ce me semble, qu'elles ne peuvent guère être déduites uniquement des expériences, et que l'expérience ne peut trancher la question de savoir si les équations de Lorenz sont préférables, par exemple, aux équations ordinaires.

Lorenz a pourtant par ses équations atteint le but de démontrer qu'il est possible d'établir des équations fondamentales applicables à tous les milieux sans recourir à l'hypothèse des forces moléculaires. C'est précisément cette hypothèse qu'il rejette et dont il fait la critique en disant (cinquième mémoire, p. 141): „On pourrait croire qu'on aurait été conduit par la connaissance si importante de la nature des forces moléculaires à des conclusions nouvelles: en réalité, il n'en fut pas ainsi. Au contraire, il s'est produit ici un fait qui s'est répété ailleurs, par exemple dans les hypothèses qu'il était nécessaire d'ajouter pour l'explication de la double réfraction: on a reconnu qu'on ne pouvait se servir des suppositions nouvelles qu'à l'endroit où elles étaient faites; on ne pouvait pas en déduire d'autres conséquences“.

On peut encore faire d'autres objections contre l'application pratique des équations fondamentales de Lorenz. Ce sont des équations aux dérivées partielles, dont l'application exige toujours une hypothèse sur la forme et la disposition des molécules; par exemple la théorie de la double réfraction implique la supposition de couches périodiques. L'admissibilité de ces hypothèses ne peut

evidemment être démontrée que par leur concordance avec la réalité; mais il est toujours difficile de conclure des effets aux causes.

De plus les applications exigent l'intégration d'équations aux dérivées partielles, dont on ne connaît pas l'intégrale générale. Les intégrations sont presque toujours faites par des développements en série, dont on suppose la forme connue *a priori*, mais de cette manière on ne sait pas si la solution du problème ainsi obtenue est la plus générale, et encore est-il souvent nécessaire d'omettre la démonstration de la convergence des séries.

Enfin, j'appellerai l'attention sur le fait qu'il est quelquefois impossible de reconnaître l'ordre de grandeur des quantités qui entrent dans les calculs de Lorenz; il me semble parfois qu'il néglige des quantités du même ordre que les quantités conservées.

Mais, quoi qu'il en soit, je crois que les équations fondamentales de Lorenz conserveront leur importance, et qu'il a pleinement démontré qu'on peut avec avantage remplacer la théorie des forces moléculaires, qui sont tout à fait hypothétiques, par la théorie plus simple de la forme et de la disposition des molécules.

Enfin j'appellerai l'attention sur le sixième mémoire „Sur l'identité des vibrations de la lumière et des courants électriques“. C'est peut-être le plus remarquable des mémoires purement théoriques de Lorenz.

Lorenz y développe dès 1867 le fondement de la théorie électromagnétique de la lumière. Il est vrai que Maxwell avait développé cette théorie deux ans auparavant, mais Lorenz n'a pas connu le travail de Maxwell et le fond des deux théories diffère. Je n'entrerai pas dans plus de détails sur le mémoire, mais je ren-

verrai le lecteur à la note 16, tome I (p. 204) où j'ai fait une comparaison approfondie des deux théories de Maxwell et de Lorenz.

Les deux mémoires suivants portent ce titre commun „Recherches expérimentales et théoriques sur les indices de réfraction“. Le premier mémoire ne contient que des expériences sur l'eau, le deuxième sur des corps différents, tant liquides que gazeux. Le but de ces deux mémoires est de trouver „l'indice réduit de réfraction“ des différents corps, c'est-à-dire l'indice de réfraction correspondant à une longueur d'onde infiniment grande. Cet indice est déterminé tant par l'expérience que par des développements purement théoriques.

Comme un rayon lumineux correspondant à une longueur d'onde infiniment grande n'existe pas dans la réalité, l'indice réduit de réfraction ne peut pas être observé directement; mais on peut le déduire des expériences en partant de la supposition que l'indice de réfraction soit exprimable par une série de la forme

$$n = A + \frac{B}{\lambda^2} + \frac{C}{\lambda^4} + \dots,$$

où λ est la longueur d'onde, A, B, C désignent des fonctions qui dépendent de la température et du volume du corps.

Ce qui importe surtout, c'est de savoir si A , l'indice réduit, dépend à la fois du volume et de la température ou du volume seul. La question n'est pas absolument tranchée par l'expérience; mais Lorenz croit pourtant probable, d'après les expériences mêmes, que l'indice est fonction du volume seul, ce qui indiquerait que les corps sont composés de molécules, dont la position seule varie

avec la température, mais qui elles-mêmes restent invariables.

Les recherches théoriques ont pour point de départ les équations aux dérivées partielles établies dans les quatrième et cinquième mémoires. Lorenz suppose de plus que les corps ordinairement appelés homogènes ne sont pas en réalité parfaitement homogènes, mais qu'ils sont composés d'éléments à structure périodique ou qui se répètent périodiquement. De plus il admet que les périodes des éléments sont petites en comparaison de la longueur d'une onde lumineuse. Par contre, il ne fait pas d'hypothèses sur la forme des molécules.

Avec ces hypothèses le calcul fait ressortir la façon dont la fonction $\frac{A^2-1}{A^2+2}v$ (la constante de réfraction), où A est l'indice réduit et v le volume du corps, dépend de l'indice de réfraction moléculaire et de la constitution du corps. Si l'on suppose le corps composé de molécules invariables, séparées par des intervalles vides, cette fonction est une constante.

Pour mener les calculs au bout, on se sert pourtant de la théorie des moyennes, théorie mal fondée au point de vue mathématique, et de plus, l'application qu'en fait Lorenz ne me semble pas toujours correcte. Les calculs sont très pénibles. Si l'on suppose les molécules sphériques, les calculs se simplifient considérablement et l'on obtient du reste les mêmes résultats. C'est ce qu'a fait voir Lorenz dans un résumé de ces deux mémoires inséré au tome XL des annales de Wiedemann, et dont les développements mathématiques figurent dans les „œuvres scientifiques“, dans un supplément aux deux mémoires. La théorie se trouve en bonne concordance avec les expériences.

XIII

L'avant-dernier mémoire est intitulé „Théorie de la dispersion“, c'est-à-dire théorie de l'indice de réfraction, considéré comme fonction de la longueur d'onde. Les développements sont ici purement mathématiques et n'empruntent rien à l'expérience; mais, tandis que la théorie de l'indice de réfraction réduit peut être développée sans aucune hypothèse sur la forme des molécules, il faut ici en faire une. Lorenz suppose que les molécules sont composées de couches sphériques; c'est pourquoi il développe d'abord les équations générales dont on a besoin pour le calcul du mouvement lumineux dans un pareil milieu, composé de couches sphériques, concentriques et homogènes. Puis il suppose encore que le mouvement lumineux s'opère de la même manière sur toute la surface d'une telle molécule, que la distance de deux molécules voisines est très petite en comparaison d'une longueur d'onde et que les molécules sont séparées par le vide. Enfin il fait cette dernière hypothèse, que l'indice de réfraction est infiniment grand dans la couche la plus proche du centre de la molécule. Les calculs sont très compliqués; mais les résultats sont relativement simples. Je ne crois pourtant pas qu'on puisse attribuer une grande importance à ces résultats, à cause de la multiplicité des hypothèses admises, dont la dernière surtout semble tout à fait arbitraire.

Le mémoire le plus important des travaux purement optiques est peut-être le dernier „Sur la lumière réfléchie et réfractée par une sphère transparente“. C'est une œuvre où les difficultés accumulées semblent dépasser les forces d'un seul homme; Lorenz l'a pourtant achevée.

Le procédé de développement est en principe le même que dans les mémoires précédents. Les équations aux dérivées partielles sont intégrées par des séries de fonctions sphériques et cylindriques; mais, tandis que dans le mémoire précédent la sommation des séries est facilitée par la supposition que les rayons des sphères, dans lesquelles a lieu le mouvement lumineux, sont petites en comparaison de la longueur des ondes lumineuses, ici les séries ne convergent que très lentement, le rayon de la sphère étant grand en comparaison de la longueur d'une onde. C'est pourquoi il faut ici recourir à une méthode qui permette de remplacer les séries par des expressions nouvelles plus simples.

La méthode de Lorenz consiste à remplacer les séries par leurs moyennes. Cette méthode est sujette à une foule d'objections. Elle ne peut être qu'approximativement juste et son admissibilité ne peut guère être prouvée en toute rigueur. On ne peut pas trouver les limites des erreurs commises en remplaçant les séries par leurs moyennes. Les séries que l'on somme de cette manière contiennent un nombre fini, mais indéterminé, de termes, et ce nombre n'est limité que par la condition qu'il doit être très grand.

De plus Lorenz conclut, de ce que la moyenne d'une série est nulle, que ses dérivées auront de même leur moyenne nulle. Mais, nonobstant ces objections, les résultats prouvent que la méthode est pratiquement applicable. Comme, en effet, elle met en évidence tous les résultats connus de la réflexion et la réfraction d'un rayon lumineux dans une sphère transparente, je crois qu'on peut conclure, avec une grande probabilité, que

les résultats nouveaux, qui peuvent être considérés comme des interpolations, sont aussi valables.

Je mentionnerai enfin une particularité des séries employées par Lorenz pour représenter le mouvement lumineux en un point donné: leurs différents termes expriment la partie du mouvement lumineux produite par un rayon dont la distance au rayon central peut être mesurée par l'indice du terme, et les coefficients des termes sont exprimés par des séries dont le m -ième terme correspond à la partie du mouvement produite par une réflexion m -uple.

Nous venons de passer en revue les travaux optiques de Lorenz qui font le contenu du premier volume. Le tome second comprend le reste des œuvres physiques de Lorenz et de plus ses travaux mathématiques. Quelques-uns des derniers travaux se rattachent pourtant à des phénomènes physiques. D'un autre côté, le premier mémoire du tome second „Mémoire sur la théorie de l'élasticité des corps homogènes à élasticité constante“ est presque exclusivement mathématique. Le problème est traité presque entièrement de la même manière que les problèmes optiques, et si j'ai renvoyé ce mémoire au tome second, c'était uniquement pour que le tome premier ne contînt que des mémoires d'optique.

Le reste des travaux physiques diffère fort des œuvres optiques. Ils traitent de la détermination des unités dites absolues, de la connexion entre les conductibilités calorifique et électrique et enfin des décharges périodiques, c'est-à-dire des phénomènes qui rattachent l'électricité et la lumière.

Le premier petit mémoire purement physique du volume „Sur le nombre des molécules contenues dans

un milligramme d'eau⁴ tend à déterminer ce nombre par la quantité d'électricité nécessaire à la décomposition d'un milligramme d'eau, par la tension minima de cette électricité et enfin par des suppositions sur la situation mutuelle des molécules. Comme résultat, Lorenz trouve que le nombre cherché doit être plus grand que $1300 \cdot 10^{18}$ et la distance des molécules plus petite que $\frac{1}{10^7}$ mm.

Lorenz prenant comme point de départ des suppositions arbitraires sur la forme et la situation des molécules, on ne peut attribuer une grande importance à ses résultats, qui indiquent plutôt l'ordre de grandeur des quantités cherchées.

Le mémoire suivant traite de la détermination du degré de chaleur en unités absolues. Tandis qu'on peut d'une manière naturelle définir les unités électriques et magnétiques en mesures absolues, il n'en est pas de même pour les unités de chaleur.

Lorenz cherche à trouver une telle mesure en prenant comme point de départ la loi de Dulong et Petit, que la quantité de chaleur nécessaire pour élever d'un même nombre de degrés la température d'un même nombre de molécules est indépendante de la nature des corps, au moins pour les gaz, proposition qui pourtant n'est pas admissible en toute rigueur. De plus Lorenz choisit un nombre déterminé de molécules. En se laissant guider par la loi de Faraday, que la même quantité d'électricité dégage des quantités équivalentes des différents électrolytes, il arrive à cette définition: Un degré de chaleur en mesures absolues est l'élévation de température que l'unité de travail peut produire, si elle est complètement transformée en chaleur, sur un nombre d'atomes d'un corps simple égal à celui que l'unité

d'électricité dégage d'un électrolyte normal (par ex. : l'acide chlorhydrique).

Le mémoire finit par quelques considérations sur la relation intime entre les conductibilités électrique et calorifique. A cause du manque de rigueur de la loi de Petit et Dulong, la détermination faite par Lorenz de l'unité absolue du degré de chaleur ne peut guère prétendre à une grande utilité pratique.

Si les deux mémoires précédents ne sont pas d'importance considérable, par contre, les trois mémoires suivants méritent la plus sérieuse considération, tant à cause de leurs idées simples qu'à cause des résultats qu'ils ont fait connaître. Ils concernent la détermination de la résistance du mercure en mesures absolues. Comme on le sait, l'unité de résistance est définie par la résistance d'un conducteur de longueur égale à l'unité par lequel passe l'unité de courant, la différence de tension électrique des deux extrémités du conducteur étant égale à l'unité. Cette unité de résistance est de la dimension d'une vitesse. Le raisonnement de Lorenz procède de cette idée, que si l'unité de résistance doit être considérée comme une vitesse, elle doit aussi être mesurée par une vitesse. En fait, Lorenz la mesure par la vitesse d'un disque tournant. Un courant passe par le corps dont on veut mesurer la résistance. Une partie du courant est dérivée et passe par le disque tournant et par une bobine de fil de cuivre qui entoure le disque. Le disque tournant induit un courant dans la bobine et le mouvement du disque est tel que le sens du courant induit est contraire à celui du courant direct. Quand le disque reçoit un mouvement convenable, les deux courants se détruisent. Mais la vitesse qu'il faut imprimer au disque

pour qu'il en soit ainsi est égale à la résistance du conducteur, multipliée par une constante qui dépend de la forme de l'appareil.

Lorenz a consacré trois mémoires à ce sujet. Dans le premier, il donne la théorie de l'appareil que nous venons de mentionner et les résultats des expériences faites avec cet appareil, qui était d'une simplicité surprenante. Ainsi Lorenz le fait tourner simplement avec la main. Dans le second mémoire, il donne une analyse des méthodes à employer pour déterminer l'Ohm et dans la troisième il expose le détail des expériences faites avec un appareil perfectionné. La comparaison des résultats des deux séries d'expériences fait ressortir combien est ingénieuse l'idée de Lorenz, car elle ne manifeste que des différences insignifiantes entre les résultats, bien que dans le premier cas l'appareil soit extrêmement simple, comme nous l'avons dit, tandis que dans l'appareil perfectionné on a mis à profit nombre de raffinements techniques.

De même le mémoire suivant „Sur la propagation de l'électricité“ présente le plus haut intérêt. Comme on le sait, Feddersen a le premier fait des expériences sur les décharges oscillantes d'une bouteille de Leyde; ensuite, Kirchhoff a appliqué la théorie à l'explication des résultats de Feddersen. Mais il s'est manifesté, à propos de la durée des oscillations, une divergence considérable entre les calculs et les observations. Ce sont ces faits que Lorenz a soumis à un examen plus précis. Il détermine d'abord, tant par le calcul que par l'observation, les constantes d'induction d'un grand nombre de fils conducteurs. Malheureusement, il n'indique que le point de départ et le résultat de ses recherches. Or

le résultat ne concorde pas avec le point de départ. J'ai pensé que cette divergence était due à une erreur dans les formules. Vraisemblablement il n'en est pas ainsi, mais la différence tient à ce que Lorenz s'est servi d'une formule prise ailleurs, et qui ne s'accorde pas avec le point de départ qu'il indique (voir la remarque de Lorenz, au bas de la page 171 du tome II). Il faut pourtant remarquer que, si Lorenz avait conservé son point de départ, ses calculs et ses expériences concorderaient presque complètement (voir, par exemple, Note 7, p. 236). Quoi qu'il en soit, le résultat final de Lorenz, que la constante diélectrique du verre est beaucoup plus grande que celle qu'avait admise Kirchhoff, explique complètement la divergence entre les calculs de Kirchhoff et les observations de Feddersen. De plus le mémoire contient une recherche des constantes d'induction du fer et de l'influence inductrice de la terre sur les fils télégraphiques. On peut remarquer que c'est Lorenz qui le premier a employé les téléphones pour étudier les effets de l'induction.

Le dernier grand mémoire du volume qui porte sur la physique concerne les conductibilités électrique et calorifique des métaux. C'est essentiellement un grand travail expérimental, où sont déterminées par deux méthodes différentes les conductibilités calorifiques des corps. Ces conductibilités calorifiques sont comparées avec les conductibilités électriques et Lorenz arrive à ce résultat qu'on peut approximativement (mais avec une assez grossière approximation) évaluer leur rapport à une constante multipliée par la température absolue. Finalement Lorenz a exécuté quelques calculs pour résoudre cette question: comment peut-on, en supposant les corps constitués d'une manière

donnée, imaginer que la conduction calorifique soit remplacée par des courants électriques locaux de façon à établir théoriquement la loi fournie par l'observation?

Malheureusement, les expériences faites au cours de ces recherches (1881) ont eu une action fâcheuse sur la santé de Lorenz, les expériences ayant été en partie exécutées dans une chambre froide.

La dernière partie du tome second des œuvres scientifiques de Lorenz contient ses travaux purement mathématiques. La plupart de ces travaux n'offrent pas un intérêt considérable. Il s'en trouve pourtant parmi eux quelques-uns d'assez remarquables. Je citerai, par exemple, le beau mémoire „Contribution à la théorie des nombres“, et je signalerai à l'attention des mathématiciens le mémoire „Sur le développement des fonctions arbitraires au moyen de fonctions données“ qui expose une théorie assez intéressante du développement des fonctions arbitraires en séries de fonctions Besséliennes. Ici les idées de Lorenz sont, je crois d'une grande portée et tout à fait originales.

D'une autre nature est l'intérêt du mémoire „Sur la compensation des erreurs d'observation“, où Lorenz cherche à prouver que si l'on a observé plusieurs fois des grandeurs O , fonctions linéaires de plusieurs autres, on n'obtiendra pas toujours la compensation la plus favorable en se servant des valeurs vraies des coefficients de ces fonctions linéaires. Il cherche en outre une méthode propre à résoudre cette question: Quelle est, parmi plusieurs compensations qui pourraient être employées, celle qui doit être considérée comme la plus favorable? Le mémoire met en évidence tout le génie de Lorenz à traiter les questions difficiles et son habileté

à combiner les raisonnements; toutefois les résultats ne me semblent que peu satisfaisants. Tout le mémoire est difficile à comprendre et manque de clarté.

Le plus étendu des mémoires purement mathématiques est le dernier „Recherches analytiques sur les nombres de nombres premiers“. Il forme une suite aux recherches antérieures sur le même sujet et cherche à démontrer que les écarts entre les nombres effectifs de nombres premiers et les nombres calculés par la formule de Riemann sont périodiques. D'après une communication que m'a faite M. Gram, il est vraisemblable que plusieurs des résultats de Lorenz sont admissibles; mais il manque à ses recherches le degré d'exactitude mathématique qui seul peut mettre les conclusions hors de doute. Les résultats obtenus sont dus plutôt au tact fin de Lorenz qu'à ses raisonnements. Les calculs sont tous approximatifs et l'on n'a aucun moyen de trouver les limites des erreurs commises.

A part ses travaux scientifiques, Lorenz a eu un grand penchant et un vif intérêt pour les applications pratiques. Ainsi il s'est occupé d'expériences nombreuses sur la production galvanique des métaux légers. En collaboration avec le professeur C.-P. Jürgensen il a construit une dynamo; le principe qui y est appliqué à des aimants intérieurs a plus tard reçu des applications d'une grande extension.

Mais c'est de la théorie et des applications du téléphone que Lorenz s'occupa surtout et, comme nous l'avons dit plus haut: il est certainement le premier qui se soit servi du téléphone pour des mesures scientifiques. Enfin Lorenz a publié nombre d'ouvrages, dont certains ont eu plusieurs éditions. Son livre „Kortfattet Naturlære“

(Abrégé de physique) a complété cinq éditions (1865—87). Un cours de physique plus étendu parut en 1870, puis vinrent „Lærebog om Lyset“ (optique) en 1876, et „Læren om Varmen“ (théorie de la chaleur), livre qui a été traduit en allemand (1877). Ces deux derniers ouvrages surtout ont de l'importance à cause de leur mode d'exposition et des points de vue pour la plupart très originaux de l'auteur. Comme on le voit, le champ d'action de Lorenz a été très étendu et son génie a enrichi la science, principalement la physique mathématique, d'une foule d'idées nouvelles.

En terminant cette longue entreprise, dont le but a été le rendre les travaux de Lorenz accessibles au monde savant, je tiens à remercier la fondation Carlsberg, qui a assumé toutes les dépenses et à la libéralité de laquelle les „œuvres scientifiques“ ont dû de voir le jour.

(Sur la vie et les œuvres de Lorenz voir: C. Christiansen, Dansk biografisk Lexikon, tome X, p. 376—381, 1896. Upsala Fyra-hundraårs jubelfest 1877, p. 363. Illustr. Tid. 1891, Nr. 38. Danskeren VI. Une partie de la présente notice sur la vie et les travaux de Lorenz est traduite de l'article de C. Christiansen.)

SUR

LE DÉVELOPPEMENT DES FONCTIONS

AU MOYEN D'INTÉGRALES DÉFINIES.

SUR LE DÉVELOPPEMENT DES FONCTIONS AU MOYEN D'INTÉGRALES DÉFINIES.

TIDSSKRIFT FOR MATEMATIK 1860, P. 160—168.*

* NOTE 1.

On connaît plusieurs méthodes qui servent à exprimer des fonctions arbitraires au moyen d'intégrales définies; ces développements ont eu surtout de l'importance pour la théorie de l'intégration des équations différentielles sous des conditions données: aussi est-ce surtout en physique mathématique qu'on a fait des applications étendues de ces méthodes. Qu'il me soit permis ici d'appeler l'attention sur une méthode nouvelle, de mettre en évidence la relation qui existe entre elle et les méthodes connues et enfin de faire ressortir une partie des applications étendues qu'on peut faire de cette méthode dans toutes les parties de la théorie de l'élasticité.

Si b est positif, on arrive aisément, au moyen de l'intégration par parties, à l'identité

$$\int_0^{\infty} e^{-bx} \cos ax \, dx = \frac{b}{a^2 + b^2}.$$

Multipliant par da et intégrant entre les limites $a = 0$ et $a = a$, on en tire

$$\int_0^{\infty} \frac{e^{-bx} \sin ax \, dx}{x} = \operatorname{arc} \operatorname{tg} \frac{a}{b},$$

d'où l'on déduira, en faisant $b = 0$,

$$\int_0^{\infty} \frac{\sin ax}{x} dx = \begin{cases} \frac{\pi}{2}, & a > 0 \\ -\frac{\pi}{2}, & a < 0. \end{cases} \quad (1)$$

Il faut pourtant remarquer que le premier membre de cette équation est indéterminé, et qu'on ne peut le déterminer qu'en tenant compte de l'origine de l'intégrale, qui est une valeur limite d'une expression déterminée, et doit être considéré comme une manière abrégée d'écrire l'intégrale

$$\left[\int_0^{\infty} \frac{e^{-bx} \sin ax}{x} dx \right]^{b=0},$$

expression à laquelle on devra recourir, si l'application de la notation abrégée conduit à des résultats indéterminés. La notation ci-dessus exprime qu'on doit évaluer b à zéro, non pas avant, mais après l'intégration; ici, comme partout dans la mathématique, quand on traite des quantités qui, à un certain endroit des calculs doivent être égalées à zéro, il importe essentiellement d'observer à quel moment cela se fait.

Si nous cherchons maintenant la valeur de la série $\sum e^{-bm} \cos am$, où la sommation est étendue à toutes les valeurs entières de m , depuis 0 jusqu'à l'infini et où l'on ne prend que la moitié du premier terme, nous obtiendrons, en remplaçant le cosinus par son expression en exponentielles imaginaires et sommant tous les termes, la valeur

$$\frac{1}{2} \cdot \frac{1 - e^{-2b}}{1 - 2e^{-b} \cos a + e^{-2b}}.$$

Pour $b = 0$ l'expression s'évanouit, à moins que $\cos a$ ne soit en même temps égal à 1. Si nous supposons la valeur de a comprise entre -2π et $+2\pi$, sans pourtant atteindre ces deux limites, a sera infiniment petit* et $\cos a = 1 - \frac{a^2}{2}$. Par conséquent on aura * NOTE 2.

$$\left[\Sigma e^{-bm} \cos am \right]^{b=0} = \left[\frac{b}{a^2 + b^2} \right]^{b=0},$$

valeur égale à l'intégrale définie considérée ci-dessus.

Si l'on multiplie par da les deux membres de cette équation et qu'on intègre de $a = 0$ jusqu'à $a = a$, on obtiendra

$$\Sigma \frac{\sin am}{m} = \begin{cases} \frac{\pi}{2} & \text{pour } 2\pi > a > 0, \\ -\frac{\pi}{2} & \text{pour } -2\pi < a < 0, \end{cases} \quad (2)$$

où le premier membre est écrit sous la forme abrégée, qui est plus commode et généralement employée, mais où la remarque faite ci-dessus est encore valable.

Si nous partons de l'équation

$$2 \int f(x) dx = \int_{\mu}^x f(a) da - \int_x^{\mu_1} f(a) da, \quad (3)$$

où la valeur de x est comprise entre une constante plus grande μ_1 et une constante plus petite μ , et si nous multiplions le second membre par

$$\pm \frac{2}{\pi} \int_0^{\infty} \frac{\sin((x-a)\beta) d\beta}{\beta} = 1,$$

où le signe supérieur est valable pour $x > a$, l'inférieur pour $x < a$, nous obtiendrons

$$2 \int f(x) dx = \frac{2}{\pi} \int_{\mu}^{\mu_1} f(a) da \int_0^{\infty} \frac{\sin(x-a)\beta}{\beta} d\beta.$$

En différentiant les deux membres par rapport à x , on obtiendra

$$f(x) = \frac{1}{\pi} \int_{\mu}^{\mu_1} f(a) da \int_0^{\infty} \cos(x-a)\beta \cdot d\beta, \quad (4)$$

$$\mu < x < \mu_1,$$

ce qui est le théorème de Fourier.

Si, de même, on multiplie le second membre de (3) par

$$\pm \frac{2}{\pi} \sum \frac{\sin(x-a)m}{m} = 1,$$

où le signe supérieur est valable pour $2\pi > x-a > 0$, l'inférieur pour $0 > x-a > -2\pi$, on obtiendra de la même manière

$$f(x) = \frac{1}{\pi} \sum \int_{-\pi}^{+\pi} f(a) \cos(m(x-a)) \cdot da, \quad \pi > x > -\pi, \quad (5)$$

où $x-a$ doit être compris entre les limites -2π et $+2\pi$, et où par conséquent x et a , qui tous deux peuvent parcourir la même série de valeurs, sont compris entre $-\pi$ et $+\pi$.

Ces deux développements au moyen d'intégrales définies sont connus et ont reçu des applications étendues. Il y a pourtant encore un troisième développement, qui peut être obtenu d'une manière analogue, à savoir en multipliant le second membre de (3) par

$$\pm \frac{2}{\pi} \left[\arctg \frac{x-a}{y-\beta} \right]^{y=\beta} = 1,$$

où le signe supérieur est valable pour $x > a$, le signe

inférieur pour $x < a$, et où l'on suppose que $y - \beta$ s'approche de zéro en parcourant une série de valeurs positives. On obtiendra alors

$$f(x) = \frac{1}{\pi} \left[\int_{\mu}^{\mu_1} f(a) da \frac{y - \beta}{(x - a)^2 + (y - \beta)^2} \right]^{y = \beta}, \quad \mu < x < \mu_1. \quad (6)$$

Si l'on supposait que $y - \beta$ fût originairement négatif, le second membre de (6) devrait être changé de signe.

Au moyen de ces trois méthodes on peut aussi développer sans difficulté des fonctions de plusieurs variables au moyen d'intégrales définies. Nous nous bornerons ici à l'application de la dernière méthode. L'équation (6) peut être écrite sous la forme plus générale

$$f(x, y) = \frac{1}{\pi} \left[\int_{\mu}^{\mu_1} f(a, y) da \frac{z - \gamma}{(x - a)^2 + (z - \gamma)^2} \right]^{z = \gamma}$$

On trouve facilement

$$\int_{-\infty}^{+\infty} \frac{d\beta}{r^3} = \frac{2}{(x - a)^2 + (z - \gamma)^2},$$

où

$$r = \sqrt{(x - a)^2 + (y - \beta)^2 + (z - \gamma)^2},$$

et, si l'on remplace le second membre de cette équation, qui figure dans l'équation ci-dessus, par le premier, on obtiendra

$$f(x, y) = \frac{1}{2\pi} \left[\int_{\mu}^{\mu_1} f(a, y) da \int_{-\infty}^{+\infty} \frac{z - \gamma}{r^3} d\beta \right]^{z = \gamma}.$$

Comme le second membre s'évanouit pour toutes valeurs finies de r , il faut seulement tenir compte des éléments pour lesquels $x = a$, $y = \beta$.

C'est pourquoi l'on peut remplacer y par β et choisir des limites arbitraires de α et β ; il faut pourtant que x et y soient compris entre elles. Par conséquent on obtiendra

$$f(x, y) = \frac{1}{2\pi} \left[\int d\alpha \int d\beta \frac{z-\gamma}{r^3} f(\alpha, \beta) \right]^{z=\gamma} \quad (7)$$

ou

$$f(x, y) = -\frac{1}{2\pi} \left[\frac{\partial}{\partial z} \int d\alpha \int d\beta \frac{f(\alpha, \beta)}{r} \right]^{z=\gamma} \quad (7')$$

Ici le point (x, y) est situé entre les limites de l'intégrale et l'on suppose que $z-\gamma$ est originairement positif. Dans le cas contraire le signe doit être renversé.

On peut remarquer que l'expression entre crochets dans l'équation (7) satisfait à l'équation différentielle

$$\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} = 0, \quad (8)$$

et qu'elle exprime, d'après la loi de Newton, l'attraction qu'exerce un plan dans la direction de la normale sur un point, si un élément du plan a la masse $f(\alpha, \beta) d\alpha d\beta$, et si l'on prend comme unité d'attraction l'attraction de l'unité de masse à l'unité de distance. L'équation (7) exprime alors que l'attraction exercée par le plan dans la direction de la normale sur un point qui s'approche infiniment de lui est égale à $2\pi f(x, y)$, c'est-à-dire à 2π fois la masse de l'élément le plus proche, divisée par la surface $dx dy$ de l'élément.

Pour des fonctions de plusieurs variables $(x_1, x_2, \dots, x_{n-1})$ on trouve plus généralement

$$f(x_1, x_2, \dots, x_{n-1}) \\ = \frac{\Gamma\left(\frac{n}{2}\right)}{\pi^{\frac{n}{2}}} \left[\int da_1 \int da_2 \dots \int da_{n-1} f(a_1, a_2, \dots, a_{n-1}) \frac{x_n - a_n}{r^n} \right]^{x_n = a_n} \quad (9)$$

où

$$r = \sqrt{(x_1 - a_1)^2 + (x_2 - a_2)^2 + \dots + (x_n - a_n)^2};$$

$x_1, x_2, \dots, x_{n-1}$ sont situés entre les limites de l'intégrale et $x_n - a_n$ est supposé positif à l'origine.

On peut facilement se convaincre de la validité de cette équation, si l'on remarque, qu'on n'a besoin que de tenir compte des éléments pour lesquels $x_1 = a_1, x_2 = a_2 \dots$. On peut donc remplacer $f(a_1, a_2, \dots, a_{n-1})$ par $f(x_1, x_2, \dots, x_{n-1})$ et remplacer les limites arbitraires par $-\infty$ et $+\infty$. Alors l'équation (9) est vérifiée sans difficulté.*

L'expression entre crochets satisfait à l'équation

$$\frac{\partial^2}{\partial x_1^2} + \frac{\partial^2}{\partial x_2^2} \dots \frac{\partial^2}{\partial x_n^2} = 0. \quad (10)$$

Par conséquent cette méthode pour exprimer les fonctions peut surtout être employée pour l'intégration des équations différentielles de la forme (10) dont l'intégrale est une fonction donnée pour une valeur donnée d'une des variables, de manière que les deux fonctions arbitraires qui entrent dans l'intégrale de (10) sont déterminées par cette fonction et par les limites de l'intégrale.

L'intégrale est donc dans ce cas déterminée par le second membre de (9), si l'on supprime les crochets.

Mais la méthode peut encore servir à l'intégration d'autres équations différentielles après avoir été légèrement généralisée. Nous prendrons comme exemple l'équation différentielle qui exprime la loi à laquelle sont assu-

* Cfr. Jacobi. Journal de Crelle, vol. 12, p. 59 et 60.

jettis tous les petits mouvements des corps élastiques à élasticité constante, savoir

$$\frac{\partial^3 \varphi}{\partial x^3} + \frac{\partial^3 \varphi}{\partial y^3} + \frac{\partial^3 \varphi}{\partial z^3} = \frac{1}{\omega^3} \frac{\partial^3 \varphi}{\partial t^3}. \quad (11)$$

On peut facilement se convaincre que cette équation est vérifiée par

$$\frac{f\left(t - \frac{r}{\omega}, \beta, \gamma\right)}{r} \quad \text{où} \quad r = \sqrt{x^2 + (y - \beta)^2 + (z - \gamma)^2};$$

par conséquent aussi par

$$\varphi = -\frac{1}{2\pi} \frac{\partial}{\partial x} \int d\beta \int d\gamma \frac{f\left(t - \frac{r}{\omega}, \beta, \gamma\right)}{r}. \quad (12)$$

Cette expression se réduit, d'après (7'), à $f(x, y, z)$ pour $x = 0$. Car, seul, le premier terme est à considérer dans le développement de $\varphi\left(t - \frac{r}{\omega}, \beta, \gamma\right)$ suivant les puissances de $\frac{r}{\omega}$, si $r = 0$, de sorte qu'on peut négliger tous les termes suivants.

Si le mouvement, qui est déterminé par l'intégrale, émane d'un plan ($x = 0$) à limites données, et si l'on suppose que l'intégrale est ici exprimée par $f(t, y, z)$ et qu'elle ne soit pas assujettie à d'autres conditions, l'équation différentielle est pour x positif complètement intégrée par (12), où l'intégrale est étendue aux limites du plan, et pour x négatif on ne doit que changer le signe.

L'intégrale (12) fait ressortir qu'on peut considérer chaque élément $d\beta d\gamma$ comme l'origine d'une partie du mouvement en (x, y, z) , qui en émane avec une vitesse

constante ω , la partie de l'intégrale due à cet élément étant

$$-\frac{1}{2\pi} d\beta d\gamma \frac{\partial}{\partial x} \frac{f\left(t - \frac{r}{\omega}, \beta, \gamma\right)}{r}.$$

Si la distance du point considéré au plan primitivement mis en mouvement est finie, on peut considérer chaque élément du plan, $d\beta d\gamma$, comme un centre de mouvement, de manière que le mouvement du point considéré devienne la somme des mouvements qui émanent des différents centres de mouvement.

C'est la proposition qui exprime le principe d'Huygens; mais ce principe est incomplet parce qu'il ne suffit pas à déterminer le mouvement. Il est de plus inexact, et si l'on s'en sert (comme le fait Lamé dans ses „Leçons“) pour le combiner avec les équations ordinaires du mouvement, on parviendra à des résultats incorrects, parce qu'il n'est pas valable aux distances infiniment petites du plan en mouvement.

Dans ce dernier cas, le mouvement doit être considéré comme émanant, non pas d'un point mais d'un élément $d\beta d\gamma$ du plan, et l'expression exacte de la partie de l'intégrale qui en provient sera, si β et γ sont les coordonnées du point milieu de l'élément

$$-\frac{1}{2\pi} \frac{\partial}{\partial x} \int_{\beta - \frac{d\beta}{2}}^{\beta + \frac{d\beta}{2}} \int_{\gamma - \frac{d\gamma}{2}}^{\gamma + \frac{d\gamma}{2}} d\beta d\gamma \frac{f\left(t - \frac{r}{\omega}, \beta, \gamma\right)}{r},$$

expression qui, d'après (7'), se réduit à $f(t, y, z)$ pour $x = 0$.

Comme exemple de l'application de l'équation (12), on peut supposer que l'intégrale puisse dans toute l'étendue d'un plan ($x = 0$) être exprimée par $f(\omega t - my - nz)$. L'intégrale peut alors être exprimée sous une forme plus simple, et le problème qui consiste à déduire cette forme de l'équation (12) ne sera pas dépourvu d'intérêt.

NOTES.

NOTE 1. On trouvera une grande partie de la substance de ce mémoire développée dans les deux premiers mémoires des tomes premier et second des œuvres scientifiques de Lorenz. Ces deux mémoires font de plus ressortir les applications qu'on peut faire des théories développées ici. Du reste la méthode dont on s'est servi ici diffère de celle qui est employée dans les mémoires cités.

NOTE 2. C'est-à-dire, si b converge vers zéro, a doit également converger vers la même limite, si l'expression ne s'évanouit pas.

UN THÉORÈME
SUR
LA FONCTION POTENTIELLE.

UN THÉORÈME SUR LA FONCTION POTENTIELLE.

MATHEMATISK TIDSSKRIFT, 1863, P. 52—63.

L'importance de la fonction potentielle pour l'application de la mathématique à la nature est tellement considérable, qu'une connaissance approfondie de ses propriétés est indispensable pour la solution d'un grand nombre de problèmes présentés par la nature elle-même.

Mais, nonobstant les recherches nombreuses faites sur cette fonction, on aura pourtant, il me semble, en traitant les problèmes de la physique mathématique, parfois le sentiment, que nos connaissances sont très limitées à son égard. La présente recherche, provoquée par différentes difficultés de la physique mathématique, qui peuvent être rapportées à la solution du problème en question, a encore confirmé ma conviction que notre connaissance de la fonction potentielle est encore restreinte, et que la présente contribution à son étude n'est que peu de chose par rapport à ce qui reste encore à faire. Que cela soit pour les autres, comme ce l'est pour moi-même, une incitation à pousser avec énergie l'étude de cette fonction!

La fonction potentielle d'un corps par rapport à un point (x, y, z) est exprimée par

$$V = \int d\alpha \int d\beta \int d\gamma \frac{f(\alpha, \beta, \gamma)}{\sqrt{(x-\alpha)^2 + (y-\beta)^2 + (z-\gamma)^2}}.$$

Les intégrations relatives aux trois variables α, β, γ peuvent ici être étendues de $-\infty$ jusqu'à $+\infty$, la fonction $f(\alpha, \beta, \gamma)$ étant tout à fait arbitraire, même discontinue, et par exemple égale à zéro dans certaines portions de l'espace. On sait alors qu'on peut inversement déterminer f par la fonction potentielle V , car on a

$$\frac{\partial^2 V}{\partial x^2} + \frac{\partial^2 V}{\partial y^2} + \frac{\partial^2 V}{\partial z^2} = -4\pi f(x, y, z).$$

Je chercherai à montrer comment le problème correspondant peut être résolu pour la fonction potentielle d'un plan par rapport à un point du plan lui-même, problème dont la solution ne peut pas être déduite de la proposition citée.

Cette fonction est, pour un plan que nous prendrons comme plan des xy , définie par

$$\varpi(x, y) = \int_{-\infty}^{+\infty} d\alpha \int_{-\infty}^{+\infty} d\beta \frac{f(\alpha, \beta)}{V(x-\alpha)^2 + (y-\beta)^2} \quad (1)$$

Le problème est d'invertir cette équation, de manière à déterminer la fonction f par la fonction ϖ .

La fonction potentielle peut être exprimée en coordonnées polaires sous la forme

$$\varpi(p, \varphi) = \int_0^\infty r dr \int_0^{2\pi} d\theta \frac{f(r, \theta)}{Vp^2 + r^2 - 2pr \cos(\theta - \varphi)}. \quad (2)$$

Avant de passer au détail de la transformation de ces intégrales, il faut commencer par établir quelques propositions auxiliaires, qui serviront à la solution du problème.

I.

On peut mettre l'intégrale

$$\int_0^{\infty} r dr \int_0^{2\pi} d\theta \frac{f(r) \cos m\theta}{\sqrt{p^2 + r^2 - 2pr \cos \theta}} \quad (3)$$

sous une autre forme, en développant le dénominateur suivant les puissances croissantes de r , si $r < p$, ou de p , si $r > p$, intégrant ensuite terme à terme, puis sommant la série ainsi obtenue (voir Ramus: Méch., p. 88). Le même résultat s'obtient par l'application de la formule (voir Schlömilch: Höh. Analyse, p. 339)*

* NOTE 1.

$$\begin{aligned} & \int_0^{2\pi} \frac{\cos m\theta d\theta}{\sqrt{1 + a^2 - 2a \cos \theta}} \\ &= a^m \int_0^{2\pi} \left(\frac{\sin \theta}{\sqrt{1 + a^2 - 2a \cos \theta}} \right)^{2m} \frac{d\theta}{\sqrt{1 + a^2 - 2a \cos \theta}}, \quad a < 1, \end{aligned}$$

où l'on a posé (voir l'endroit cité)

$$\frac{\sin \theta}{\sqrt{1 + a^2 - 2a \cos \theta}} = \sin \phi,$$

d'où résulte

$$\frac{1}{\sqrt{1 - a^2 \sin^2 \phi}} = \frac{\sqrt{1 + a^2 - 2a \cos \theta}}{1 - a \cos \theta}, \quad d\phi = \frac{1 - a \cos \theta}{1 + a^2 - 2a \cos \theta} d\theta.$$

En multipliant ces deux équations, on obtiendra

$$\frac{d\phi}{\sqrt{1 - a^2 \sin^2 \phi}} = \frac{d\theta}{\sqrt{1 + a^2 - 2a \cos \theta}},$$

ce qui transforme l'intégrale ci-dessus en

$$a^m \int_0^{2\pi} \frac{\sin^{2m} \phi \, d\phi}{\sqrt{1 - a^2 \sin^2 \phi}},$$

les limites étant les mêmes pour ϕ que pour θ , car $\sin \phi$ doit prendre toutes les valeurs entre -1 et $+1$, limites qui correspondent à $\cos \theta = a$.

Comme $\sin \phi$ n'entre dans l'intégrale que par son carré, on peut encore faire varier ϕ de 0 à $\frac{\pi}{2}$, et quadrupler le résultat, ce qui donne

$$\int_0^{2\pi} \frac{\cos m\theta \, d\theta}{\sqrt{1 + a^2 - 2a \cos \theta}} = 4 a^m \int_0^{\frac{\pi}{2}} \frac{\sin^{2m} \phi \, d\phi}{\sqrt{1 - a^2 \sin^2 \phi}}.$$

Si l'on remplace ici a par $\frac{r}{p}$ ou par $\frac{p}{r}$ suivant que $r \leq p$, on reconnaîtra aisément que l'intégrale (3) peut être exprimée par

$$4 \int_0^{\frac{\pi}{2}} \sin^{2m} \phi \, d\phi \left[\int_0^p \frac{r \, dr \, f(r)}{\sqrt{p^2 - r^2 \sin^2 \phi}} \frac{r^m}{p^m} + \int_p^\infty \frac{r \, dr \, f(r)}{\sqrt{r^2 - p^2 \sin^2 \phi}} \frac{p^m}{r^m} \right]. \quad (4)$$

Dans la première de ces deux intégrales, je pose $r \sin \phi = p \sin \varphi$, par où j'obtiens

$$\frac{d\phi}{\sqrt{p^2 - r^2 \sin^2 \phi}} = \frac{d\varphi}{\sqrt{r^2 - p^2 \sin^2 \varphi}}.$$

Or l'intégration peut être exécutée de deux manières, soit d'abord par rapport à φ , les limites de $\sin \varphi$ étant 0 et $\frac{r}{p}$, puis par rapport à r entre les limites $r = 0$ et $r = p$, soit d'abord par rapport à r , entre les limites $r = p \sin \varphi$ et $r = p$, puis par rapport à φ de $\varphi = 0$ à $\varphi = \frac{\pi}{2}$. Dans les deux cas les deux variables par-

courront la même série de valeurs; mais on doit ici préférer la dernière méthode, parce qu'elle permet d'exécuter immédiatement l'addition des deux termes de (4), en remplaçant de nouveau φ par ψ . Cette expression sera alors transformée en

$$4 \int_0^{\frac{\pi}{2}} \sin^{2m} \psi \, d\psi \int_{p \sin \psi}^{\infty} \frac{r \, dr \, f(r)}{\sqrt{r^2 - p^2 \sin^2 \psi}} \frac{p^m}{r^m}.$$

Si l'on pose ici

$$p \sin \psi = q, \quad d\psi = \frac{dq}{\sqrt{p^2 - q^2}},$$

on obtiendra par conséquent

$$\begin{aligned} & \int_0^{\infty} r \, dr \int_0^{2\pi} d\theta \frac{f(r) \cos m\theta}{\sqrt{p^2 + r^2 - 2pr \cos \theta}} \\ &= 4 \int_0^p \frac{dq}{\sqrt{p^2 - q^2}} \int_q^{\infty} \frac{r \, dr \, f(r)}{\sqrt{r^2 - q^2}} \frac{q^{2m}}{r^m p^m}. \end{aligned} \quad (5)$$

On pourrait transformer la seconde intégrale de (4) d'une manière analogue, en posant

$$p \sin \psi = r \sin \varphi,$$

d'où

$$\frac{d\psi}{\sqrt{r^2 - p^2 \sin^2 \psi}} = \frac{d\varphi}{\sqrt{p^2 - r^2 \sin^2 \varphi}}.$$

Comme précédemment, on suppose que l'intégration est exécutée en premier lieu par rapport à r entre les limites $r = p$ et $r = \frac{p}{\sin \varphi}$ et puis par rapport à φ de $\varphi = 0$ à $\varphi = \frac{\pi}{2}$.

Si l'on remplace de nouveau φ par ψ et si l'on fait l'addition des deux termes de (4), on trouvera

$$4 \int_0^{\frac{\pi}{2}} \sin^{2m} \phi \, d\phi \int_0^{\frac{p}{\sin \phi}} \frac{r \, dr \, f(r)}{V p^2 - r^2 \sin^2 \phi} \frac{r^m}{p^m},$$

et, si l'on fait ici la substitution $\frac{p}{\sin \phi} = q$, on obtiendra une nouvelle expression de l'intégrale en question (3), savoir

$$\begin{aligned} & \int_0^\infty r \, dr \int_0^{2\pi} d\theta \frac{f(r) \cos m \theta}{V p^2 + r^2 - 2pr \cos \theta} \\ &= 4 \int_p^\infty \frac{dq}{V q^2 - p^2} \int_0^q \frac{r \, dr \, f(r)}{V q^2 - r^2} \frac{r^m p^m}{q^{2m}}. \end{aligned} \quad (6)$$

On reconnaît pourtant facilement que les deux expressions (5) et (6) sont identiques, ce qui devient évident si l'on pose dans la première $q = pu$ et $r = puv$, et dans la seconde $q = pv$ et $r = pvu$.

II.

Abel a démontré qu'une fonction d'une seule variable peut être exprimée au moyen d'une intégrale définie de la manière suivante (voir Ramus: Diff. et Int., p. 269).*

$$\varphi(x) = \frac{\sin m\pi}{\pi} \int_0^x \frac{d\theta}{(x-\theta)^{1-m}} \int_0^\theta \frac{\varphi'(z) \, dz}{(\theta-z)^m}, \quad (7)$$

où m est une fraction proprement dite, positive et arbitraire.

Nous poserons ici

$$m = \frac{1}{2}, \quad \theta = q^2, \quad z = r^2, \quad x = p^2,$$

puis nous remplacerons $\varphi(p^2)$ par $\varphi(p)$. On obtiendra

ainsi

$$\varphi(p) = \frac{2}{\pi} \int_0^p \frac{q dq}{V p^2 - q^2} \int_0^q \frac{\varphi'(r) dr}{V q^2 - r^2}. \quad (8)$$

On peut se servir de cette expression si l'on a, par exemple,

$$\varphi(p) = \frac{2}{\pi} \int_0^p \frac{q dq}{V p^2 - q^2} \phi(q). \quad (9)$$

pour déterminer inversement la fonction ϕ par φ , car on doit avoir

$$\phi(q) = \int_0^q \frac{\varphi'(r) dr}{V q^2 - r^2}. \quad (9')$$

Au contraire, si dans la formule (7) nous posons

$$m = \frac{1}{2}, \quad \theta = \frac{1}{q^2}, \quad z = \frac{1}{r^2}, \quad x = \frac{1}{p^2},$$

et si nous remplaçons ensuite $\varphi\left(\frac{1}{p^2}\right)$ par $\varphi(p)$, nous obtiendrons

$$\varphi(p) = -\frac{2}{\pi} \int_p^\infty \frac{p dq}{q V q^2 - p^2} \int_q^\infty \frac{r \varphi'(r) dr}{V r^2 - q^2}. \quad (10)$$

Cette équation va être différenciée par rapport à p ; mais il faut observer qu'on doit alors, pour éviter des indéterminations provenant de termes infinis, exécuter la différenciation en supposant d'abord que la limite inférieure soit $p+h$, puis intégrer par parties et enfin faire tendre h vers zéro*. Si l'on écrit de nouveau $\varphi(p)$ à la place de $\varphi'(p)$, on obtiendra aisément de cette manière

$$\varphi(p) = -\frac{2}{\pi} \int_p^\infty \frac{dq}{V q^2 - p^2} \cdot \frac{d}{dq} \int_q^\infty \frac{r \varphi(r) dr}{V r^2 - q^2}. \quad (11)$$

* NOTE 3.

La différentiation indiquée par rapport à q peut être exécutée de la même manière, par où l'équation prendra la forme

$$\varphi(p) = -\frac{2}{\pi} \int_p^\infty \frac{q \, dq}{Vq^2 - p^2} \int_q^\infty \frac{\varphi'(r) \, dr}{Vr^2 - q^2}. \quad (12)$$

III.

A présent nous passerons de ces recherches préliminaires au problème en question. Il sera alors naturel de chercher en premier lieu la solution d'un cas particulier, celui où la fonction $f(r, \theta)$ de l'équation (2) ne dépend que de r seulement. On peut dans ce cas remplacer $\theta - \varphi$ dans l'équation (2) par θ , tandis que les limites de θ seront constamment 0 et 2π , et $\varpi(p, \varphi)$ sera alors fonction de p seulement.

Nous pouvons alors partir de l'équation

$$\varpi(p) = \int_0^\infty r \, dr \int_0^{2\pi} d\theta \frac{f(r)}{Vp^2 + r^2 - 2pr \cos \theta}, \quad (13)$$

qui, au moyen de la relation (5), où $m = 0$, sera transformée en

$$\varpi(p) = 4 \int_0^p \frac{dq}{Vp^2 - q^2} \int_q^\infty \frac{r f(r) \, dr}{Vr^2 - q^2}.$$

Si nous posons ici

$$\int_q^\infty \frac{r f(r) \, dr}{Vr^2 - q^2} = q \phi(q),$$

nous aurons

$$\varpi(p) = 4 \int_0^p \frac{q \, dq \phi(q)}{Vp^2 - q^2},$$

équation qui est identique à l'équation (9), si l'on pose $\frac{\varpi(p)}{2\pi} = \varphi(p)$. L'équation suivante (9') deviendra alors

$$\phi(q) = \frac{1}{2\pi} \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}},$$

d'où, en introduisant l'expression trouvée ci-dessus pour $\phi(q)$, l'on déduira

$$\frac{1}{q} \int_0^\infty \frac{r f(r) dr}{\sqrt{r^2 - q^2}} = \frac{1}{2\pi} \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}}.$$

Si maintenant on multiplie cette équation par $-\frac{2}{\pi} \frac{p dq}{\sqrt{q^2 - p^2}}$ et qu'on intègre entre les limites $q = p$ et $q = \infty$, on obtiendra

$$-\frac{2}{\pi} \int_p^\infty \frac{p dq}{\sqrt{q^2 - p^2}} \int_q^\infty \frac{r f(r) dr}{\sqrt{r^2 - q^2}} = -\frac{1}{\pi^2} \int_p^\infty \frac{p dq}{\sqrt{q^2 - p^2}} \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}}.$$

Le premier membre est d'après (10) égal à $\int f(p) dp$ si l'on remplace $\varphi'(r)$ par $f(r)$, et l'on tirera de l'équation, en différentiant par rapport à p ,

$$f(p) = -\frac{1}{\pi^2} \frac{d}{dp} \int_p^\infty \frac{p dq}{\sqrt{q^2 - p^2}} \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}}.$$

Le problème est donc résolu, car la fonction f est exprimée au moyen de la fonction ϖ , mais il y a avantage à mettre le résultat sous d'autres formes.

Si la différentiation par rapport à p est exécutée de la manière indiquée ci-dessus dans un cas analogue, on obtiendra *

$$f(p) = -\frac{1}{\pi^2} \int_p^\infty \frac{dq}{\sqrt{q^2 - p^2}} \frac{d}{dq} \left(q \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}} \right),$$

puis

$$f(p) = -\frac{1}{\pi^2} \int_p^\infty \frac{dq}{\sqrt{q^2 - p^2}} \int_0^q \frac{dr}{\sqrt{q^2 - r^2}} \frac{d(r\varpi'(r))}{dr}. \quad (14)$$

Sous cette forme le résultat se prête mieux au calcul pratique. Du reste il est identique au résultat suivant

NOTE 5.
$$f(p) = -\frac{1}{\pi^2} \frac{1}{p} \frac{d}{dp} \left(p \frac{d}{dp} \int_p^\infty \frac{dq}{\sqrt{q^2 - p^2}} \int_0^q \frac{dr r \varpi(r)}{\sqrt{q^2 - r^2}} \right)^*, \quad (15)$$

ce qu'on peut reconnaître en effectuant les différentiations indiquées.

Si nous cherchons au contraire à donner le résultat sous une forme analogue à l'équation (13), nous pourrions nous servir de l'équation (6), où nous ferons $m = 0$ et d'abord $f(r) = \frac{d(r\varpi'(r))}{r dr}$, puis $f(r) = \varpi(r)$. Par là les équations (14) et (15) se transformeront en

$$f(p) = -\frac{1}{4\pi^2} \int_0^\infty r dr \int_0^{2\pi} \frac{d\theta}{\sqrt{p^2 + r^2 - 2pr \cos \theta}} \frac{d(r\varpi'(r))}{r dr} \quad (16)$$

et

$$f(p) = -\frac{1}{4\pi^2} \frac{1}{p} \frac{d}{dp} \left(p \frac{d}{dp} \int_0^\infty r dr \int_0^{2\pi} \frac{d\theta \cdot \varpi(r)}{\sqrt{p^2 + r^2 - 2pr \cos \theta}} \right), \quad (17)$$

de manière à représenter $f(p)$ sous forme d'une nouvelle fonction potentielle.

Après avoir résolu le problème dans le cas spécial considéré, nous passerons au cas général. Soit proposée l'équation

$$\varpi(p, \varphi) = \int_0^\infty r dr \int_0^{2\pi} d\theta \frac{f(r, \theta)}{\sqrt{p^2 + r^2 - 2pr \cos(\theta - \varphi)}},$$

il est à prévoir que la fonction $f(r, \theta)$ pourra être exprimée sous une forme analogue au résultat obtenu récemment; c'est pourquoi nous examinerons en détail l'expression

$$F = \int_0^\infty r dr \int_0^{2\pi} d\theta \frac{\varpi(r, \theta)}{\sqrt{p^2 + r^2 - 2pr \cos(\theta - \varphi)}}. \quad (18)$$

Ici on a

$$\varpi(r, \theta) = \int_0^\infty r' dr' \int_0^{2\pi} d\theta' \frac{f(r', \theta')}{\sqrt{r^2 + r'^2 - 2rr' \cos(\theta' - \theta)}} \quad (19)$$

où nous exprimerons $f(r', \theta')$, d'une manière bien connue, par une intégrale définie (voir Ramus: Diff. et Int., p. 82), savoir

$$f(r', \theta') = \frac{1}{\pi} \sum \int_{-\pi}^{+\pi} f(r', \psi) \cos m(\theta' - \psi) d\psi. \quad (20)$$

Σ désigne ici une somme étendue à toutes les valeurs entières et positives de m , le premier terme devant être divisé par 2. Le second membre de (18) se présente par ces substitutions sous la forme d'une intégrale quintuple, où nous remplacerons θ par $\theta + \varphi$ et θ' par $\theta + \theta' + \varphi$, sans pourtant changer les limites de l'intégrale; $\cos m(\theta' - \psi)$ sera donc remplacé par $\cos m(\theta' + \theta + \varphi - \psi)$, mais, comme les termes qui contiennent $\sin m\theta$ et $\sin m\theta'$ s'évanouiront dans l'intégration, on peut remplacer $\cos m(\theta' + \theta + \varphi - \psi)$ par

$$\cos m\theta \cos m\theta' \cos m(\varphi - \psi).$$

Si l'on pose, pour abrégér,

$$\phi(r) = \int_0^\infty r' dr' \int_0^{2\pi} d\theta' \frac{f(r', \phi) \cos m\theta'}{Vr'^2 + r'^2 - 2r'r' \cos \theta'},$$

F sera exprimée par

$$F = \frac{1}{\pi} \sum \int_{-\pi}^{+\pi} d\phi \cos m(\varphi - \phi) \int_0^\infty r dr \int_0^{2\pi} d\theta \frac{\phi(r) \cos m\theta}{Vp^2 + r^2 - 2pr \cos \theta}.$$

Ici nous pouvons employer les formules trouvées antérieurement, car nous aurons d'après (5)

$$\phi(r) = 4 \int_0^r \frac{dq'}{Vr'^2 - q'^2} \int_{q'}^\infty \frac{r' dr' f(r', \phi)}{Vr'^2 - q'^2} \frac{q'^{2m}}{r'^m r^m},$$

et d'après (6)

$$\int_0^\infty r dr \int_0^{2\pi} d\theta \frac{\phi(r) \cos m\theta}{Vp^2 + r^2 - 2pr \cos \theta} = 4 \int_p^\infty \frac{dq}{Vq^2 - p^2} \int_0^q \frac{r dr \phi(r)}{Vq^2 - r^2} \frac{r^m p^m}{q^{2m}},$$

et par conséquent l'expression de F sera

$$\begin{aligned} & \frac{16}{\pi} \sum \int_{-\pi}^{+\pi} d\phi \cos m(\varphi - \phi) \int_p^\infty \frac{dq}{Vq^2 - p^2} \\ & \times \int_0^q \frac{r dr}{Vq^2 - r^2} \int_0^r \frac{dq'}{Vr'^2 - q'^2} \int_{q'}^\infty \frac{r' dr' f(r', \phi)}{Vr'^2 - q'^2} \frac{q'^{2m} p^m}{q^{2m} r^m}. \end{aligned}$$

Mais, d'après (8), on a

$$\int_0^q \frac{r dr}{Vq^2 - r^2} \int_0^r \frac{dq' \phi'(q')}{Vr'^2 - q'^2} = \frac{\pi}{2} \phi(q).$$

Par suite, l'expression ci-dessus se transformera en

$$F = 8 \sum \int_{-\pi}^{+\pi} d\phi \cos m(\varphi - \phi) \int_p^\infty \frac{dq \phi(q)}{Vq^2 - p^2} \frac{p^m}{q^{2m}}.$$

Si nous différencions par rapport à p , nous obtenons

$$p \frac{\partial F}{\partial p} = 8 \sum \int_{-\pi}^{+\pi} d\psi \cos m(\varphi - \psi) \int_p^{\infty} \frac{q dq}{\sqrt{q^2 - p^2}} [q \varphi'(q) - m \varphi(q)] \frac{p^m}{q^{2m+1}}.$$

Comme on a de plus

$$\frac{d}{dp} \int_p^{\infty} \frac{q dq}{\sqrt{q^2 - p^2}} \frac{\varphi'(q)}{q^{2m}} = \int_p^{\infty} \frac{p dq}{\sqrt{q^2 - p^2}} \frac{d}{dq} \cdot \frac{\varphi'(q)}{q^{2m}}$$

et

$$\frac{d}{dp} \int_p^{\infty} \frac{dq}{\sqrt{q^2 - p^2}} \cdot \frac{\varphi(q)}{q^{2m}} = \int_p^{\infty} \frac{q dq}{p \sqrt{q^2 - p^2}} [q \varphi'(q) - 2m \varphi(q)] \frac{1}{q^{2m+1}},$$

si l'on différencie de nouveau l'expression ci-dessus par rapport à p , on aura

$$\begin{aligned} & \frac{1}{p} \frac{\partial}{\partial p} \left(p \frac{\partial F}{\partial p} \right) \\ = & 8 \sum \int_{-\pi}^{+\pi} d\psi \cos m(\varphi - \psi) \int_p^{\infty} \frac{dq}{\sqrt{q^2 - p^2}} \left[p^m \frac{d}{dq} \frac{\varphi'(q)}{q^{2m}} + m^2 p^{m-2} \frac{\varphi(q)}{q^{2m}} \right]. \end{aligned}$$

Si l'on remarque que la dernière partie de l'intégrale est égale à $-\frac{1}{p^2} \frac{\partial^2 F}{\partial \varphi^2}$, en faisant passer ce terme au premier membre de l'équation, on obtiendra

$$\begin{aligned} & \frac{1}{p} \frac{\partial}{\partial p} \left(p \frac{\partial F}{\partial p} \right) + \frac{1}{p^2} \frac{\partial^2 F}{\partial \varphi^2} \\ = & 8 \sum \int_{-\pi}^{+\pi} d\psi \cos m(\varphi - \psi) p^m \int_p^{\infty} \frac{dq}{\sqrt{q^2 - p^2}} \frac{d}{dq} \left(\frac{\varphi'(q)}{q^{2m}} \right). \quad (21) \end{aligned}$$

Ici il faut substituer, d'après l'expression trouvée ci-dessus de $\varphi'(q')$,

$$\frac{\varphi'(q)}{q^{2m}} = \int_q^\infty \frac{r' dr' f(r', \psi)}{V_{r'^2 - q^2}} \frac{1}{r'^m},$$

et nous pourrons nous servir de l'équation (11), si nous remplaçons r par r' , $\varphi(r)$ par $\frac{f(r', \psi)}{r'^m}$, pour réduire le second membre de l'équation (21), qui deviendra

$$-4\pi \sum \int_{-\pi}^{+\pi} d\psi \cos m(\varphi - \psi) f(p, \psi).$$

Mais, d'après (20), cette expression est identique à

$$-4\pi^2 f(p, \varphi),$$

de manière que l'équation (21) se réduit à

$$\frac{1}{p} \frac{\partial}{\partial p} \left(p \frac{\partial F'}{\partial p} \right) + \frac{1}{p^2} \frac{\partial^2 F'}{\partial \varphi^2} = -4\pi^2 f(p, \varphi). \quad (22)$$

Le problème est donc résolu, car ici $f(p, \varphi)$ est exprimé, au moyen de F , et si l'on introduit la valeur (18) de F , on aura

$$f(p, \varphi) = -\frac{1}{4\pi^2} \left(\frac{1}{p} \frac{\partial}{\partial p} \cdot p \frac{\partial}{\partial p} + \frac{1}{p^2} \frac{\partial^2}{\partial \varphi^2} \right) \int_0^\infty r dr \int_0^{2\pi} \frac{\varpi(r, \theta) d\theta}{V_{p^2 + r^2 - 2pr \cos(\theta - \varphi)}}. \quad (23)$$

En introduisant les coordonnées rectangulaires, on tirera de l'équation de départ (1), ou

$$\varpi(a, \beta) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{f(a', \beta')}{V_{(a-a')^2 + (\beta-\beta')^2}}, \quad (24)$$

d'après le résultat trouvé

$$^* \text{ NOTE 6. } f(x, y) = -\frac{1}{4\pi^2} \left(\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} \right) \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{\varpi(a, \beta)}{V_{(x-a)^2 + (y-\beta)^2}}. \quad (25)$$

On peut écrire le second membre de cette équation sous une forme plus commode pour le calcul, en remplaçant α par $\alpha+x$, β par $\beta+y$ puis effectuant les différentiations indiquées et enfin remplaçant de nouveau $\alpha+x$ par α , $\beta+y$ par β , par où l'on obtiendra

$$f(x, y) = -\frac{1}{4\pi^2} \int_{-\infty}^{+\infty} d\alpha \int_{-\infty}^{+\infty} \frac{d\beta}{\sqrt{(x-\alpha)^2 + (y-\beta)^2}} \left(\frac{\partial^2}{\partial \alpha^2} + \frac{\partial^2}{\partial \beta^2} \right) \varpi(\alpha, \beta). \quad (26)$$

D'une manière analogue on transformera (23) en

$$f(p, \varphi) = -\frac{1}{4\pi^2} \int_0^\infty r dr \int_0^{2\pi} \frac{d\theta}{\sqrt{p^2 + r^2 - 2pr \cos(\theta - \varphi)}} \left(\frac{1}{r} \frac{\partial}{\partial r} r \frac{\partial}{\partial r} + \frac{1}{r} \frac{\partial^2}{\partial \theta^2} \right) \varpi(r, \theta). \quad (27)$$

NOTES.

NOTE 1. On a

$$\begin{aligned} u &= \int_0^{2\pi} \frac{\cos m\theta d\theta}{\sqrt{1+a^2-2a\cos\theta}} = 2 \int_0^{\pi} \frac{\cos m\theta d\theta}{\sqrt{1+a^2-2a\cos\theta}} \\ &= \frac{2a}{m} \int_0^{\pi} \frac{\sin m\theta \sin \theta d\theta}{(\sqrt{1+a^2-2a\cos\theta})^3}. \end{aligned}$$

Mais, d'après un théorème bien connu de Jacobi, on aura

$$\begin{aligned} &\sin(m \arccos x) \\ &= (-1)^{m-1} \frac{d^{m-1}(1-x^2)^{m-\frac{1}{2}}}{dx^{m-1}} \cdot \frac{m}{1 \cdot 3 \cdot 5 \dots (2m-1)}, \end{aligned}$$

et par conséquent, en remplaçant dans l'expression ci-dessus $\cos \theta$ par x ,

$$u = \frac{2a(-1)^{m-1}}{1 \cdot 3 \cdot 5 \dots (2m-1)} \int_{-1}^{+1} \frac{d^{m-1}(1-x^2)^{m-\frac{1}{2}}}{dx^{m-1}} \frac{1}{(\sqrt{1+a^2-2ax})^3} dx.$$

En intégrant ici $(m-1)$ fois par parties, on obtiendra

$$u = 2a^m \frac{^{+1}_{-1} (1-x^2)^{m-\frac{1}{2}} dx}{(\sqrt{1+a^2-2ax})^{2m+1}},$$

ce qui est précisément l'expression de Lorenz, si l'on remplace de nouveau x par $\cos \theta$.

NOTE 2. Comme toute la suite de ce mémoire est fondée sur le théorème d'Abel (Oeuvres complètes, tome I, p. 27), je développerai ici la démonstration d'Abel, légèrement modifiée. On a, en remplaçant z par θz , θ par θx ,

$$\int_0^x \frac{\theta^s d\theta}{(x-\theta)^{1-m}} \int_0^\theta \frac{z^{a-1} dz}{(\theta-z)^m} = x^{a+s} \int_0^1 \frac{\theta^{a-m+s} d\theta}{(1-\theta)^{1-m}} \int_0^1 \frac{z^{a-1} dz}{(1-z)^m}.$$

Si l'on suppose ici $0 < m < 1$, $a - m + s > -1$, $a > 0$, on aura

$$\int_0^1 \frac{\theta^{a-m+s} d\theta}{(1-\theta)^{1-m}} = \frac{\Gamma(a-m+s+1) \Gamma(m)}{\Gamma(a+s+1)},$$

$$\int_0^1 \frac{z^{a-1} dz}{(1-z)^m} = \frac{\Gamma(a) \Gamma(1-m)}{\Gamma(a-m+1)},$$

et par suite, comme

$$\Gamma(m) \Gamma(1-m) = \frac{\pi}{\sin m\pi},$$

$$x^{a+s} = \frac{\sin m\pi}{\pi} \frac{\Gamma(a+s+1) \Gamma(a-m+1)}{\Gamma(a) \Gamma(a-m+s+1)} \int_0^x \frac{\theta^s d\theta}{(x-\theta)^{1-m}} \int_0^\theta \frac{z^{a-1} dz}{(\theta-z)^m}.$$

Si l'on pose ici $s = 0$, on aura

$$x^a = \frac{\sin m\pi}{\pi} \int_0^x \frac{d\theta}{(x-\theta)^{1-m}} \int_0^\theta \frac{az^{a-1} dz}{(\theta-z)^m}.$$

En multipliant les deux membres de cette équation par $f(a) da$ et intégrant de $a = a$ à $a = b$, on aura

$$\int_a^b f(a) x^a da = \frac{\sin m\pi}{\pi} \int_0^x \frac{d\theta}{(x-\theta)^{1-m}} \int_0^\theta \frac{a f(a) z^{a-1} da}{(\theta-z)^m} dz.$$

Si l'on pose enfin $\int_a^b f(a) x^a da = \varphi(x)$, on obtiendra

$$(7) \quad \varphi(x) = \frac{\sin m\pi}{\pi} \int_0^\theta \frac{d\theta}{(x-\theta)^{1-m}} \int_0^\theta \frac{\varphi'(z) dz}{(\theta-z)^m}.$$

Mais ici se présente cette question: La fonction $\varphi(x)$ peut-elle être tout à fait arbitraire?

On reconnaît immédiatement qu'elle doit s'évanouir pour $x = 0$ et qu'elle doit être différentiable. Puis on reconnaît que, la formule (7) étant valable pour chaque terme d'une somme, elle l'est encore pour toute la somme. En conséquence, si $\varphi'(z)$ peut être développé en une série de puissances à exposants positifs, la série représentera $f(x) - f(0)$. On peut encore conclure de là que la série représentera $f(x) - f(0)$, si $f'(x)$ peut être développé en une série de Fourier.

Si l'on pose $s = -1$, on obtiendra

$$x^{a-1} = \frac{\sin m\pi}{\pi} (a-m) \int_0^x \frac{d\theta}{\theta(x-\theta)^{1-m}} \int_0^\theta \frac{z^{a-1} dz}{(\theta-z)^m},$$

formule qui n'est valable que pour $a > m$.

On en déduit

$$x^{a-m} = \frac{\sin m\pi}{\pi} \int_0^x \frac{d\theta \cdot x^{1-m}}{\theta(x-\theta)^{1-m}} \int_0^\theta \frac{(a-m) z^{a-m-1} \cdot z^m dz}{(\theta-z)^m},$$

puis, multipliant par $f(a) da$, $f(a)$ étant arbitraire, et intégrant entre les limites a et b ,

$$\begin{aligned} & \int_a^b f(a) x^{a-m} da \\ &= \frac{\sin m\pi}{\pi} \int_0^x \frac{d\theta x^{1-m}}{\theta(x-\theta)^{1-m}} \int_0^\theta \frac{z^m \left(\int_a^b f(a) (a-m) z^{a-m-1} da \right) dz}{(\theta-z)^m}. \end{aligned}$$

Si l'on pose $\int_a^b f(a) x^{\alpha-m} da = \varphi(x)$, on aura

$$\varphi(x) = \frac{\sin m\pi}{\pi} \int_0^x \frac{x^{1-m} d\theta}{\theta(x-\theta)^{1-m}} \int_0^\theta \frac{\varphi'(z) z^m}{(\theta-z)^m} dz.$$

Comme ci-dessus, $\varphi(x)$ n'est pas tout à fait arbitraire, mais doit s'évanouir pour $x = 0$. Si l'on remplace ici x, θ, z par $x^{-\frac{1}{m}}, \theta^{-\frac{1}{m}}, z^{-\frac{1}{m}}$ et $\varphi(x^{-\frac{1}{m}})$ par $\varphi(x)$, on obtiendra

$$\varphi(x) = -\frac{\sin m\pi}{m\pi} \int_x^\infty \frac{\theta^{\frac{1-m}{m}} d\theta}{\left(\theta^{\frac{1}{m}} - x^{\frac{1}{m}}\right)^{1-m}} \int_\theta^\infty \frac{\varphi'(z) dz}{\left(z^{\frac{1}{m}} - \theta^{\frac{1}{m}}\right)^m}.$$

Pour $m = \frac{1}{2}$, on trouve la formule (12) de Lorenz. Ici $\varphi(x)$ doit s'évanouir pour $x = \infty$ et, du reste, on peut analyser la validité de cette formule comme celle de la formule (7). Si $\varphi(\infty)$ n'est pas égal à zéro, on doit remplacer $\varphi(x)$ par $\varphi(x) - \varphi(\infty)$.

NOTE 3. Au lieu de procéder comme le fait Lorenz, on peut se servir d'une autre méthode, qui me semble plus facile. On a

$$\varphi(p) = -\frac{2}{\pi} \int_p^\infty \frac{p dq}{q \sqrt{q^2 - p^2}} \int_q^\infty \frac{r \varphi'(r) dr}{r \sqrt{r^2 - q^2}}.$$

Si l'on pose

$$\psi(q) = \int_q^\infty \frac{r \varphi'(r) dr}{r \sqrt{r^2 - q^2}},$$

et si l'on remplace q par pq , on aura

$$\varphi(p) = -\frac{2}{\pi} \int_1^\infty \frac{dq}{q \sqrt{q^2 - 1}} \psi(pq),$$

$$\begin{aligned}\varphi'(p) &= -\frac{2}{\pi} \int_1^\infty \frac{dq}{Vq^2-1} \varphi'(pq) \\ &= -\frac{2}{\pi} \int_p^\infty \frac{dq}{Vq^2-p^2} \varphi'(q).\end{aligned}$$

Si l'on remplace ici φ' par φ , on obtiendra la formule (11).

Par la même méthode on obtiendra

$$\varphi'(q) = \frac{d}{dq} \int_q^\infty \frac{r\varphi(r) dr}{Vr^2-q^2} = \frac{1}{q} \int_q^\infty \frac{r \frac{d(r\varphi(r))}{dr}}{Vr^2-q^2} dr.$$

Mais on a

$$\begin{aligned}\frac{1}{q} \int_q^\infty \frac{r \frac{d(r\varphi(r))}{dr}}{Vr^2-q^2} dr &= \frac{1}{q} \int_q^\infty \frac{r^2 \varphi'(r) + r\varphi(r)}{Vr^2-q^2} dr \\ &= q \int_q^\infty \frac{\varphi'(r) dr}{Vr^2-q^2} + \frac{1}{q} \lim_{r=\infty} V\overline{r^2-q^2} \varphi(r).\end{aligned}$$

Si $\lim_{r=\infty} V\overline{r^2-q^2} \varphi(r) = 0$, on obtiendra l'expression de Lorenz.

NOTE 4. Les différentiations peuvent être exécutées comme dans la note 3.

NOTE 5. Dans la note 3, nous avons trouvé que

$$\begin{aligned}\frac{d \int_p^\infty \frac{q\psi(q) dq}{Vq^2-p^2}}{dp} &= \frac{1}{p} \int_p^\infty \frac{q \frac{d(q\psi(q))}{dq}}{Vq^2-p^2} dq \\ &= p \int_p^\infty \frac{\psi'(q)}{Vq^2-p^2} dq,\end{aligned}\tag{I}$$

si $\left| \psi(q) V\overline{q^2-p^2} \right|_p^\infty = 0$.

De la même manière on obtient

$$\begin{aligned} \frac{d \int_0^q \frac{r \varpi(r)}{\sqrt{q^2 - r^2}} dr}{dq} &= \frac{1}{q} \int_0^q \frac{d(r \varpi(r))}{\sqrt{q^2 - r^2}} dr \\ &= q \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}}. \end{aligned} \quad (\text{II})$$

Si l'on se sert de ces deux formules, on obtiendra

$$\frac{d}{dp} \int_p^\infty \frac{dq}{\sqrt{q^2 - p^2}} \int_0^q \frac{r \varpi(r)}{\sqrt{q^2 - r^2}} dq = \frac{1}{p} \int_p^\infty \frac{q dq}{\sqrt{q^2 - p^2}} \frac{d}{dq} \int_0^q \frac{r \varpi(r)}{\sqrt{q^2 - r^2}} dr,$$

si dans la formule (I) on pose

$$\phi(q) = \frac{1}{q} \int_0^q \frac{r \varpi(r)}{\sqrt{q^2 - r^2}} dr,$$

et si l'on se sert ensuite de la formule (II), l'expression prendra la forme

$$\frac{1}{p} \int_p^\infty \frac{q^2 dq}{\sqrt{q^2 - p^2}} \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}}.$$

En exécutant sur cette expression l'opération $\frac{1}{p} \frac{d}{dp} \cdot p$ on obtiendra

$$\frac{1}{p} \frac{d}{dp} \int_p^\infty \frac{q^2 dq}{\sqrt{q^2 - p^2}} \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}},$$

et, si dans la formule (I) on pose

$$\phi(q) = q \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}},$$

on trouvera

$$\int_p^\infty \frac{dq}{\sqrt{q^2 - p^2}} \frac{d}{dq} q \int_0^q \frac{\varpi'(r) dr}{\sqrt{q^2 - r^2}},$$

où l'on suppose que

$$\left| q \sqrt{q^2 - p^2} \int_0^q \frac{\varpi'(r) dr}{V \sqrt{q^2 - r^2}} \right|_p^\infty = 0.$$

Si l'on se sert de nouveau de la formule (II), on retrouvera l'expression (14).

NOTE 6. On peut démontrer la proposition de Lorenz d'une manière différente et plus facile.

Si l'on a donné la densité superficielle $f(x, y)$ de chaque point (x, y) du plan $z = 0$, le potentiel d'un point arbitraire du plan sera

$$\varpi(x, y) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{d\alpha d\beta \cdot f(\alpha, \beta)}{V(x-\alpha)^2 + (y-\beta)^2},$$

et de même le potentiel d'un point arbitraire de l'espace sera

$$\Pi(x, y) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{d\alpha d\beta f(\alpha, \beta)}{V(x-\alpha)^2 + (y-\beta)^2 + z^2}.$$

Mais en supposant que

$$\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \int_0^{\infty} \left[\left(\frac{\partial}{\partial x} \right)^2 + \left(\frac{\partial}{\partial y} \right)^2 + \left(\frac{\partial}{\partial z} \right)^2 \right] \Pi dx dy dz$$

soit une quantité finie, le théorème de Lord Kelvin met en évidence qu'il n'existe qu'une seule fonction satisfaisant, comme Π , entre les limites de l'intégrale à l'équation $\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} + \frac{\partial^2 u}{\partial z^2} = 0$, et prenant pour $z = 0$ les mêmes valeurs $\varpi(x, y)$ que Π . Or la fonction

$$-\frac{1}{2\pi} \frac{\partial}{\partial z} \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{\varpi(a, \beta) da d\beta}{V(x-a)^2 + (y-\beta)^2 + z^2}$$

satisfait à ces conditions. Par conséquent, on aura

$$\begin{aligned} & \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{da d\beta f(a, \beta)}{V(x-a)^2 + (y-\beta)^2 + z^2} \\ &= -\frac{1}{2\pi} \frac{\partial}{\partial z} \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{\varpi(a, \beta) da d\beta}{V(x-a)^2 + (y-\beta)^2 + z^2} \end{aligned}$$

et comme

$$\begin{aligned} & \left| -\frac{1}{2\pi} \frac{\partial}{\partial z} \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{da d\beta f(a, \beta)}{V(x-a)^2 + (y-\beta)^2 + z^2} \right|_{z=0} = f(x, y), \\ f(x, y) &= \left| \frac{1}{4\pi^2} \frac{\partial^2}{\partial z^2} \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{\varpi(a, \beta) da d\beta}{V(x-a)^2 + (y-\beta)^2 + z^2} \right|_{z=0}. \end{aligned}$$

Mais comme

$$\frac{\partial^2 \Pi}{\partial x^2} + \frac{\partial^2 \Pi}{\partial y^2} + \frac{\partial^2 \Pi}{\partial z^2} = 0,$$

on aura

$$f(x, y) = -\frac{1}{4\pi^2} \left(\frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} \right) \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \frac{da d\beta \varpi(a, \beta)}{V(x-a)^2 + (y-\beta)^2}.$$

SUR L'ÉVALUATION DES AIRES.

SUR L'ÉVALUATION DES AIRES.

MATHEMATISK TIDSSKRIFT 1864, P. 33—38.

L'évaluation de l'aire d'une surface peut être ramenée à l'évaluation du volume d'un corps infiniment mince, de même forme que la surface, car le volume divisé par l'épaisseur du corps, qui est supposée partout la même, donnera l'aire de la surface. Comme cette conception peut offrir des avantages pratiques, ainsi que le montre, à mon avis, l'exemple suivant, j'espère qu'il ne sera dépourvu d'intérêt d'en faire une étude plus détaillée.

L'équation de la surface donnée est

$$u = F(x, y, z) = 0.$$

$x+h$, $y+k$, $z+l$ étant les coordonnées d'un point situé infiniment près du point x, y, z sur la normale à la surface en ce point, on aura

$$F(x+h, y+k, z+l) = \frac{\partial u}{\partial x}h + \frac{\partial u}{\partial y}k + \frac{\partial u}{\partial z}l,$$

et de plus

$$\frac{h}{\frac{\partial u}{\partial x}} = \frac{k}{\frac{\partial u}{\partial y}} = \frac{l}{\frac{\partial u}{\partial z}} = \frac{\varepsilon}{v},$$

où

$$\varepsilon = \sqrt{h^2 + k^2 + l^2},$$

et

$$v = \sqrt{\left(\frac{\partial u}{\partial x}\right)^2 + \left(\frac{\partial u}{\partial y}\right)^2 + \left(\frac{\partial u}{\partial z}\right)^2}.$$

ε est la distance du point à la surface.

Au moyen des valeurs de h, k, l , déterminées de cette manière, on obtient

$$F(x+h, y+k, z+l) = v\varepsilon.$$

Si l'on se déplace de cette manière à une petite distance de la surface, dans la direction de la normale à un nouveau point x, y, z , la valeur de la fonction $\frac{u}{v}$ variera de 0 à ε , et, pour un point situé sur cette normale, à une plus petite distance de la surface, la valeur de $\frac{u}{v}$ sera comprise entre 0 et ε .

Pour chaque point x, y, z appartenant au corps infiniment mince ayant même forme que la surface et une épaisseur constante ε , on aura donc

$$\varepsilon > \frac{u}{v} > 0.$$

Le volume d'un corps peut être obtenu par sommation de tous les éléments de l'espace illimité, multipliés par un facteur qui est égal à 1 pour tous les éléments appartenant au corps et égal à zéro pour tous les éléments situés en dehors du corps. Comme facteur de cette nature, on peut prendre, par exemple, l'intégrale définie

$$\frac{1}{\pi} \int_0^\varepsilon \frac{da \cdot h}{\left(\frac{u}{v} - a\right)^2 + h^2} = \frac{1}{\pi} \left[\text{arc tg } \frac{\varepsilon v - u}{hv} + \text{arc tg } \frac{u}{hv} \right],$$

si l'on fait, après l'intégration, converger h vers zéro; car on aura dans ce cas

$$\operatorname{arc\,tg} \frac{\varepsilon v - u}{h v} = \begin{cases} \frac{\pi}{2} & \text{pour } \varepsilon > \frac{u}{v} \\ -\frac{\pi}{2} & \text{pour } \varepsilon < \frac{u}{v}, \end{cases}$$

et

$$\operatorname{arc\,tg} \frac{u}{h v} = \begin{cases} \frac{\pi}{2} & \text{pour } \frac{u}{v} > 0 \\ -\frac{\pi}{2} & \text{pour } \frac{u}{v} < 0. \end{cases}$$

Les deux termes sont donc en même temps égaux à $\frac{\pi}{2}$, si

$$\varepsilon > \frac{u}{v} > 0;$$

au contraire ils seront de signes contraires pour toutes les autres valeurs de $\frac{u}{v}$, et par suite ils se détruiront.

En conséquence, on aura

$$\left[\frac{1}{\pi} \int_0^\varepsilon \frac{d\alpha v^2 h}{(u - \alpha v)^2 + v^2 h^2} \right]_{h=0} = \begin{cases} 1 \\ 0 \end{cases}$$

selon que le point x, y, z est situé à l'intérieur du corps ou à l'extérieur.

Le volume du corps se trouve en multipliant cette intégrale définie par $dx dy dz$ et intégrant dans tout l'espace illimité; l'aire de la surface est donc ce volume divisé par l'épaisseur ε du corps. Par conséquent on aura

$$A = \frac{1}{\varepsilon \pi} \int_{-\infty}^{+\infty} dx \int_{-\infty}^{+\infty} dy \int_{-\infty}^{+\infty} dz \int_0^\varepsilon \frac{d\alpha \cdot v^2 h}{(u - \alpha v)^2 + v^2 h^2}.$$

Mais, comme ε est une quantité infiniment petite et comme on a, en général, pour une fonction arbitraire $\phi(\varepsilon)$,

$$\lim_{\varepsilon} \frac{\phi(\varepsilon) - \phi(0)}{\varepsilon} = \phi'(0),$$

on aura

$$A = \frac{1}{\pi} \int_{-\infty}^{+\infty} dx \int_{-\infty}^{+\infty} dy \int_{-\infty}^{+\infty} dz \frac{v^2 h}{u^2 + v^2 h^2}. \quad (1)$$

Il va sans dire, qu'on ne doit pas intégrer par rapport à x entre $-\infty$ et $+\infty$, mais seulement entre les limites qui correspondent à deux plans donnés, dans le cas où l'on ne cherche pas l'aire totale de la surface, mais bien l'aire limitée par deux plans, que nous supposons parallèles au plan des yz .

Dans l'application de la formule (1) à l'évaluation d'une aire, on ne doit pas perdre de vue que h est une quantité infiniment petite et que, par suite, tous les éléments de l'intégrale s'évanouissent, à moins que u soit en même temps infiniment petit. C'est pourquoi l'on peut, soit conserver la forme des coefficients de h et h^2 comme fonctions de trois variables, soit exprimer ces coefficients en fonction de deux variables en éliminant la troisième variable au moyen de l'équation $u = 0$.

La formule trouvée pour A peut aisément être ramenée à la forme ordinaire, si nous multiplions le numérateur et le dénominateur par $\frac{\partial u}{\partial z}$ et si nous posons

$$A = \frac{1}{\pi} \int_{-\infty}^{+\infty} dx \int_{-\infty}^{+\infty} dy \int_{-\infty}^{+\infty} dz \frac{[v^2] \frac{\partial u}{\partial z} h}{\left[\frac{\partial u}{\partial z} \right] (u^2 + [v^2] h^2)},$$

où les crochets $[]$ indiquent que z est éliminée au moyen de l'équation $u = 0$.

En intégrant par rapport à z , nous obtiendrons

$$\int dz \frac{[v] \frac{\partial u}{\partial z} h}{u^2 + [v^2] h^2} = \text{arc tg} \frac{u}{[v] h},$$

intégrale qui est égale à π toutes les fois que u prend la valeur $u = 0$, z parcourant toutes les valeurs de $-\infty$ à $+\infty$. Par conséquent, on aura

$$A = \int dx \int dy \frac{[v]}{\left[\frac{\partial u}{\partial z} \right]}. \quad (2)$$

en même temps que $u = 0$, de manière que les limites de l'intégrale deviendront les valeurs extrêmes de x et y , admissibles sous cette condition.

Comme exemple de l'application de la formule (1), nous prendrons l'évaluation de l'aire d'un ellipsoïde.

Nous aurons ici

$$u = \frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} - 1 = 0,$$

et, par suite,

$$v^2 = 4 \left[\frac{x^2}{a^4} + \frac{y^2}{b^4} + \frac{z^2}{c^4} \right].$$

Si nous posons

$$\frac{x}{a} = r \cos \theta, \quad \frac{y}{b} = r \sin \theta \cos \omega, \quad \frac{z}{c} = r \sin \theta \sin \omega,$$

nous obtiendrons

$$dx dy dz = a^2 b^2 c^2 r^2 \sin \theta d\theta dr d\omega,$$

de manière que l'aire de la surface totale de l'ellipsoïde deviendra, l'intégrale étant étendue à tout l'espace illimité,

$$A = \frac{a^2 b^2 c^2}{\pi} \int_0^\infty r^2 dr \int_0^\pi \sin \theta d\theta \int_0^{2\pi} d\omega \frac{4r^2 h}{(a^2 r^2 - 1)^2 + 4r^2 h^2},$$

où

$$a^2 = a^2 \cos^2 \theta + b^2 \sin^2 \theta \cos^2 \omega + c^2 \sin^2 \theta \sin^2 \omega.$$

Comme tous les éléments de cette intégrale s'évalouissent à moins que $a^2 r^2 - 1 = 0$ et par suite $r = \frac{1}{a}$, on peut substituer cette valeur dans les coefficients de h et de h^2 , et l'on pourra de plus poser

$$a^2 r^2 - 1 = (ar + 1)(ar - 1) = 2(ar - 1).$$

De cette manière on obtiendra

$$A = \frac{a^2 b^2 c^2}{\pi} \int_0^\infty dr \int_0^\pi \sin \theta d\theta \int_0^{2\pi} d\omega \frac{h}{a^2 (a^2 (ar - 1)^2 + h^2)},$$

formule qui, après l'intégration par rapport à r , se réduira à

$$A = a^2 b^2 c^2 \int_0^\pi \sin \theta d\theta \int_0^{2\pi} \frac{d\omega}{a^4},$$

Si l'on introduit ici la valeur de a indiquée ci-dessus et si l'on effectue l'intégration d'abord par parties par rapport à θ et puis complètement par rapport à ω , on parviendra sans difficulté aux résultats connus. Le résultat peut s'obtenir plus facilement sous une autre forme de la manière suivante.

On a

$$\frac{\partial \frac{1}{a^2}}{\partial (a^2)} + \frac{\partial \frac{1}{a^2}}{\partial (b^2)} + \frac{\partial \frac{1}{a^2}}{\partial (c^2)} = -\frac{1}{a^4},$$

et par suite

$$A = -\frac{a^2 b^3 c^2}{2} \left(\frac{\partial}{a \partial a} + \frac{\partial}{b \partial b} + \frac{\partial}{c \partial c} \right) \int_0^\pi \sin \theta d\theta \int_0^{2\pi} \frac{d\omega}{a^2}.$$

En intégrant d'abord par rapport à ω , et appliquant la formule

$$\int_0^{2\pi} \frac{d\omega}{m^2 \cos^2 \omega + n^2} = \frac{2\pi}{n \sqrt{m^2 + n^2}},$$

on obtiendra

$$\int_0^\pi \sin \theta d\theta \int_0^{2\pi} \frac{d\omega}{a^2} = 2\pi \int_0^\pi \frac{\sin \theta d\theta}{\sqrt{a^2 \cos^2 \theta + b^2 \sin^2 \theta} \sqrt{a^2 \cos^2 \theta + c^2 \sin^2 \theta}},$$

intégrale qui peut être immédiatement mise sous forme d'intégrale elliptique par la substitution

$$\cos \theta = \frac{c}{\sqrt{a^2 - c^2}} \operatorname{tg} \phi,$$

où nous supposons $a > b > c$. L'intégrale ci-dessus se réduira ainsi à

$$4\pi \int_0^\mu \frac{d\phi}{\sqrt{b^3(a^2 - c^2) - a^2(b^2 - c^2) \sin^2 \phi}},$$

où $\mu = \arccos \frac{c}{a}$. On aura donc

$$A = -2\pi a^2 b^3 c^2 \left(\frac{\partial}{a \partial a} + \frac{\partial}{b \partial b} + \frac{\partial}{c \partial c} \right) \int_0^\mu \frac{d\phi}{\sqrt{b^3(a^2 - c^2) - a^2(b^2 - c^2) \sin^2 \phi}},$$

ce qui peut, avec la notation ordinaire des intégrales elliptiques

$$F(k, \mu) = \int_0^\mu \frac{d\phi}{\sqrt{1 - k^2 \sin^2 \phi}}, \quad k = \frac{a}{b} \sqrt{\frac{b^2 - c^2}{a^2 - c^2}},$$

s'écrire de la façon suivante

$$A = -2\pi a^2 b^3 c^2 \left(\frac{\partial}{a \partial a} + \frac{\partial}{b \partial b} + \frac{\partial}{c \partial c} \right) \frac{F(k, \mu)}{b \sqrt{a^2 - c^2}}.$$

SUR
LE MOUVEMENT PERMANENT
D'UN LIQUIDE.

SUR LE MOUVEMENT PERMANENT D'UN LIQUIDE.

TIDSSKRIFT FOR MATEMATIK 1866, P. 65—74.

Quand le mouvement d'un liquide parfait est devenu permanent, la vitesse du liquide est en chaque point constante en direction et en grandeur, et toutes les particules du liquide qui passent par le même endroit doivent avoir le même mouvement et décrire la même courbe. Il est bien connu que la pression p exercée sur chaque particule d'une telle courbe peut être déterminée, à savoir par l'équation

$$p = -\rho g z - \frac{1}{2} \rho h^2 + C,$$

où ρ est la densité constante du liquide, g la pesanteur, qui est supposée la seule force accélératrice, agissant dans la direction négative de l'axe des z et h la vitesse. Si la constante C est éliminée au moyen des valeurs correspondantes p_0 , h_0 , z_0 d'un point donné de la courbe, l'équation peut s'écrire

$$p - p_0 = \rho g (z - z_0) - \frac{1}{2} \rho (h^2 - h_0^2). \quad (1)$$

Mais tandis que la pression est de cette manière déterminée pour tous points de la même courbe, décrite par une particule du liquide, l'équation (1) n'apprend rien sur la pression en général d'un point du liquide, car les constantes peuvent varier d'une manière par-

faitement inconnue, dans le passage d'une courbe à une autre.

Nous ne pouvons, par exemple, rien conclure de cette équation, pour la pression d'un liquide qu'on a fait tourner autour d'un axe fixe, et la supposition que toutes les constantes fussent les mêmes pour toutes les courbes amènerait à des résultats faux.

Au contraire, l'équation a, comme on sait, été employée avantageusement à la détermination de la pression en tout point d'un liquide qui, par une ouverture très petite s'écoule d'un vase, et c'est pourquoi l'on doit rechercher dans quels cas il est permis de faire une telle application de l'équation.

Les mouvements permanents peuvent être divisés en deux espèces, à savoir les mouvements en courbes fermées et en courbes infinies. Les premiers ont lieu dans le cas d'un mouvement permanent proprement dit, produit en un espace fini et limité; car chaque particule du liquide doit ici, pendant la durée infinie du mouvement, finir par repasser en un point où elle est déjà passée antérieurement et, par suite, décrire une courbe fermée. Il est évident qu'on ne peut pas, dans ce cas, employer l'équation (1) à la détermination de la pression en général.

Dans le second cas, où les espaces parcourus deviennent infinis, ce qui suppose le liquide illimité, le mouvement peut, soit s'étendre à toutes les particules, et dans ce cas nous ne savons encore rien de la loi d'après laquelle la pression varie d'une courbe à une autre, soit être limité par la condition de n'atteindre à une grandeur appréciable qu'à l'intérieur d'une portion limitée

de l'espace, tandis qu'il s'évanouit au voisinage des limites de ce domaine.

Dans le dernier cas, la pression peut être déterminée en général; car, dans ce cas, chaque particule du liquide qui participe au mouvement à dû, dans la suite infinie des temps, passer en un point de l'espace où le mouvement était infiniment petit; aussi pouvons-nous, dans l'équation (1), évaluer h_0 à zéro, tandis qu'on, a d'après la loi hydrostatique,

$$p_0 = \varpi - \rho g z_0,$$

où ϖ est la pression extérieure sur le liquide, le plan des xy passant par la surface du liquide dans le même endroit.* Nous supposons, dans ce qui suit, que cette * NOTE 1. pression est constante.

On aura donc

$$p - \varpi = -\rho g z - \frac{1}{2}\rho h^2, \quad (2)$$

équation qui ne contient plus de constante susceptible de varier d'une courbe à l'autre et par laquelle la pression est, en conséquence, déterminée pour tous les points du liquide. Il va sans dire que la surface libre et toutes les surfaces de niveau sont déterminées par cette équation, c'est-à-dire par la vitesse h .

On peut évidemment faire une application approximative de ce résultat dans des cas très nombreux, savoir dans tous les cas où le mouvement d'un liquide n'est constant que pendant un certain temps et où, en outre, les courbes décrites par les particules peuvent être considérées comme ayant pris naissance à des endroits où le mouvement est très faible, ce qui, par exemple, a lieu pour un liquide qui s'écoule d'un grand vase par

un petit orifice placé au fond. Il faut pourtant que le mouvement du liquide soit négligeable à une distance suffisamment grande de cet orifice, et que par conséquent le mouvement du liquide ne soit pas un mouvement de rotation.

De plus les conditions peuvent être réalisées en d'autres cas, comme par les mouvements des corps parfaitement liquides dans des canaux, dans des fontaines, etc. Mais il va sans dire que, si l'on veut appliquer la théorie aux liquides réels, on doit, selon la nature du problème, tenir plus ou moins de compte du frottement.

A présent nous pouvons avec sûreté faire un pas en avant, en retenant les conditions de validité de l'équation trouvée.

Si l'on pose pour abrégér

$$-\frac{1}{\rho}(p - \varpi) - gz = q,$$

l'équation (2) se réduira à

$$q = \frac{1}{2}h^2 = \frac{1}{2}(u^2 + v^2 + w^2), \quad (3)$$

où u , v , w sont les composantes de la vitesse suivant les trois axes. Les équations générales de l'hydrodynamique, dans le cas des liquides en mouvement permanent, peuvent, avec les notations introduites ici, s'écrire

$$\left. \begin{aligned} \frac{\partial q}{\partial x} &= u \frac{\partial u}{\partial x} + v \frac{\partial u}{\partial y} + w \frac{\partial u}{\partial z}, \\ \frac{\partial q}{\partial y} &= u \frac{\partial v}{\partial x} + v \frac{\partial v}{\partial y} + w \frac{\partial v}{\partial z}, \\ \frac{\partial q}{\partial z} &= u \frac{\partial w}{\partial x} + v \frac{\partial w}{\partial y} + w \frac{\partial w}{\partial z}. \end{aligned} \right\} \quad (4)$$

Si, de même, on déduit de (3) les dérivées de q , on obtiendra par comparaison

$$v \frac{\partial v}{\partial x} + w \frac{\partial w}{\partial x} = v \frac{\partial u}{\partial y} + w \frac{\partial u}{\partial z}$$

etc., par où l'on parvient à ces deux équations

$$\frac{\frac{\partial v}{\partial z} - \frac{\partial w}{\partial y}}{u} = \frac{\frac{\partial w}{\partial x} - \frac{\partial u}{\partial z}}{v} = \frac{\frac{\partial u}{\partial y} - \frac{\partial v}{\partial x}}{w}. \quad (5)$$

On a de plus l'équation de continuité

$$\frac{\partial u}{\partial x} + \frac{\partial v}{\partial y} + \frac{\partial w}{\partial z} = 0. \quad (6)$$

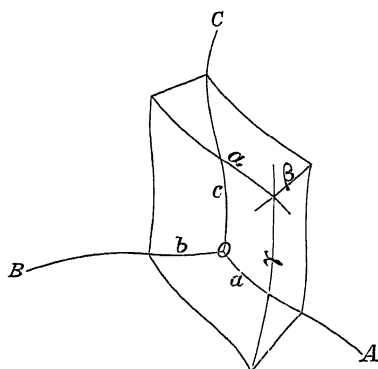
On peut remarquer le cas particulier où les numérateurs des équations (5) peuvent s'évanouir et où, par suite, u , v , w deviennent les dérivées de la même fonction par rapport à x , y , z , cas qui a été considéré à plusieurs reprises. Il est bon de remarquer que c'est précisément le cas d'un liquide qui s'écoule d'un vase par une ouverture circulaire, sous les conditions énoncées ci-dessus; car on trouve facilement dans ce cas $\frac{\partial u}{\partial y} - \frac{\partial v}{\partial z} = 0^*$, et par conséquent les autres numérateurs des équations (5) doivent aussi s'évanouir.

Pour mieux mettre en évidence le sens des équations, nous introduirons un nouveau système de coordonnées, composé de trois systèmes de trajectoires orthogonales, qui par conséquent se coupent partout sous des angles droits.

Soient pour axes coordonnés les lignes d'intersection de trois surfaces fixes, savoir OA , OB , OC , O étant l'origine; la position d'un point est déterminée par les

* NOTE 1.

trois segments a, b, c détachés sur les trois axes par les surfaces trajectoires qui passent au point considéré.



Soient α, β, γ les arcs des courbes d'intersection des trois surfaces trajectoires, comptés du point aux surfaces fixes coordonnées. Leurs prolongements, les éléments $da, d\beta, d\gamma$, font entre eux des angles droits, et soient l_1, l_2, l_3 les cosinus des angles qu'ils font avec l'axe

des x d'un système plan et rectangulaire de coordonnées; m_1, m_2, m_3 ceux des angles qu'ils font avec l'axe des y ; n_1, n_2, n_3 ceux relatifs à l'axe des z . C'est ce que représente le tableau ci-dessous

	$da \ d\beta \ d\gamma$		
x	l_1	l_2	l_3
y	m_1	m_2	m_3
z	n_1	n_2	n_3

où, comme on sait, la somme des carrés des trois quantités qui se trouvent dans une rangée horizontale ou verticale est l'unité, tandis que la somme des produits de deux quantités voisines dans deux rangées parallèles s'évanouit.

On aura donc, si a, b, c sont choisis comme variables indépendantes et si φ désigne une fonction quelconque de a, b, c ,

$$\frac{\partial \varphi}{\partial x} = \frac{\partial \varphi}{\partial a} \frac{\partial a}{\partial x} + \frac{\partial \varphi}{\partial b} \frac{\partial b}{\partial x} + \frac{\partial \varphi}{\partial c} \frac{\partial c}{\partial x}.$$

Si l'on pose ensuite

$$\frac{\partial \alpha}{\partial a} = \varepsilon_1, \quad \frac{\partial b}{\partial \beta} = \varepsilon_2, \quad \frac{\partial c}{\partial \gamma} = \varepsilon_3^*, \quad * \text{ NOTE 3.}$$

on aura

$$\frac{\partial \alpha}{\partial x} = \varepsilon_1 \frac{\partial \alpha}{\partial x} = \varepsilon_1 l_1, \quad \frac{\partial b}{\partial x} = \varepsilon_2 l_2, \quad \frac{\partial c}{\partial x} = \varepsilon_3 l_3.$$

Par conséquent, on aura

$$\left. \begin{aligned} \frac{\partial \varphi}{\partial x} &= \varepsilon_1 l_1 \frac{\partial \varphi}{\partial a} + \varepsilon_2 l_2 \frac{\partial \varphi}{\partial b} + \varepsilon_3 l_3 \frac{\partial \varphi}{\partial c}, \\ \frac{\partial \varphi}{\partial y} &= \varepsilon_1 m_1 \frac{\partial \varphi}{\partial a} + \varepsilon_2 m_2 \frac{\partial \varphi}{\partial b} + \varepsilon_3 m_3 \frac{\partial \varphi}{\partial c}, \\ \frac{\partial \varphi}{\partial z} &= \varepsilon_1 n_1 \frac{\partial \varphi}{\partial a} + \varepsilon_2 n_2 \frac{\partial \varphi}{\partial b} + \varepsilon_3 n_3 \frac{\partial \varphi}{\partial c}, \end{aligned} \right\} \quad (7)$$

où les deux dernières équations sont formées par symétrie. Si l'on multiplie ces équations respectivement par l_1 , m_1 , n_1 , et qu'on les additionne, on obtiendra la première des équations suivantes:

$$\left. \begin{aligned} \varepsilon_1 \frac{\partial \varphi}{\partial a} &= l_1 \frac{\partial \varphi}{\partial x} + m_1 \frac{\partial \varphi}{\partial y} + n_1 \frac{\partial \varphi}{\partial z}, \\ \varepsilon_2 \frac{\partial \varphi}{\partial b} &= l_2 \frac{\partial \varphi}{\partial x} + m_2 \frac{\partial \varphi}{\partial y} + n_2 \frac{\partial \varphi}{\partial z}, \\ \varepsilon_3 \frac{\partial \varphi}{\partial c} &= l_3 \frac{\partial \varphi}{\partial x} + m_3 \frac{\partial \varphi}{\partial y} + n_3 \frac{\partial \varphi}{\partial z}; \end{aligned} \right\} \quad (8)$$

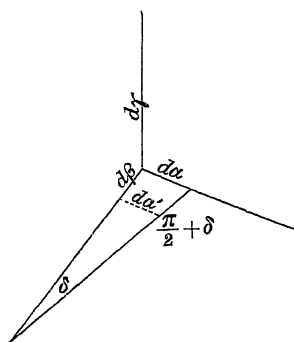
les deux dernières s'obtiennent d'une manière analogue.

Du reste, on peut encore déduire immédiatement ces trois équations.

Il faut encore se servir du développement suivant. Comme les trois systèmes se coupent suivant leurs lignes de courbure, ainsi que l'a démontré Lamé, les deux droites qui, par exemple, sont menées normalement à la surface $\alpha\gamma$ aux deux extrémités de l'élément da se

couperont en un point. Si δ est l'angle infiniment petit formé par ces deux lignes, $\frac{\delta}{da}$ sera la courbure de l'arc a

* NOTE 4.



dans la surface ay .* Tandisque l'une des deux droites fait un angle droit avec da , l'autre fera l'angle $\frac{\pi}{2} + \delta$ avec le prolongement de cet élément et, avec les axes rectangulaires, des angles dont les cosinus sont $l_2 + \frac{\partial l_2}{\partial a} da$, $m_2 + \frac{\partial m_2}{\partial a} da$, $n_2 + \frac{\partial n_2}{\partial a} da$.

Par conséquent on aura

$$\cos\left(\frac{\pi}{2} + \delta\right) = l_2\left(l_2 + \frac{\partial l_2}{\partial a} da\right) + m_2\left(m_2 + \frac{\partial m_2}{\partial a} da\right) + n_2\left(n_2 + \frac{\partial n_2}{\partial a} da\right),$$

et, par suite, comme δ est infiniment petit et

$$l_2 l_2 + m_2 m_2 + n_2 n_2 = 0,$$

$$\frac{\delta}{da} = -\left(l_2 \frac{\partial l_2}{\partial a} + m_2 \frac{\partial m_2}{\partial a} + n_2 \frac{\partial n_2}{\partial a}\right),$$

ou bien

$$\frac{\delta}{da} = \varepsilon_1 \left(l_2 \frac{\partial l_1}{\partial a} + m_2 \frac{\partial m_1}{\partial a} + n_2 \frac{\partial n_1}{\partial a}\right). \quad (9)$$

Mais cette courbure peut encore être déterminée d'une autre manière; car on trouve facilement par la figure

$$\delta = \frac{da - da'}{d\beta} = -\frac{d(da)}{d\beta}.$$

Si l'on introduit ici les valeurs

$$da = \frac{da}{\varepsilon_1} \text{ et } d\beta = \frac{db}{\varepsilon_2},$$

on trouvera

$$\partial = \frac{\varepsilon_2}{\varepsilon_1} \frac{d\varepsilon_1}{db} da,$$

et, par conséquent, la courbure est déterminée par

$$\frac{\partial}{da} = \frac{\varepsilon_2}{\varepsilon_1} \frac{\partial \varepsilon_1}{\partial b}. \quad (10)$$

Les deux équations (9) et (10) donnent donc

$$l_2 \frac{\partial l_1}{\partial a} + m_2 \frac{\partial m_1}{\partial a} + n_2 \frac{\partial n_1}{\partial a} = \frac{\varepsilon_2}{\varepsilon_1} \frac{\partial \varepsilon_1}{\partial b}. \quad (11)$$

Ces considérations finies, nous reviendrons à nos équations hydrodynamiques (4), et, supposant que les courbes α soient les chemins parcourus par les particules du liquide*, nous poserons

* NOTE 5.

$$u = hl_1, \quad v = hm_1, \quad w = hn_1.$$

On doit ici remarquer que les courbes α ne déterminent pas seulement les surfaces $\beta\gamma$, qui partout font des angles droits avec ces courbes, mais encore les deux autres systèmes de surfaces, car ces deux systèmes doivent passer par les lignes de courbure des surfaces $\beta\gamma$. D'après (7) on a en général

$$u \frac{\partial \varphi}{\partial x} + v \frac{\partial \varphi}{\partial y} + w \frac{\partial \varphi}{\partial z} = h\varepsilon_1 \frac{\partial \varphi}{\partial a},$$

de manière que les équations (4) se transforment en

$$\frac{\partial q}{\partial x} = h\varepsilon_1 \frac{\partial(hl_1)}{\partial a}, \quad \frac{\partial q}{\partial y} = h\varepsilon_1 \frac{\partial(hm_1)}{\partial a}, \quad \frac{\partial q}{\partial z} = h\varepsilon_1 \frac{\partial(hn_1)}{\partial a},$$

d'où l'on déduira au moyen des équations (8) et des relations connues entre les quantités l, m, n

$$\left. \begin{aligned} \varepsilon_1 \frac{\partial q}{\partial a} &= h \varepsilon_1 \frac{\partial h}{\partial a}, \\ \varepsilon_2 \frac{\partial q}{\partial b} &= h^2 \varepsilon_1 \left(l_2 \frac{\partial l_1}{\partial a} + m_2 \frac{\partial m_1}{\partial a} + n_2 \frac{\partial n_1}{\partial a} \right), \\ \varepsilon_3 \frac{\partial q}{\partial c} &= h^2 \varepsilon_1 \left(l_3 \frac{\partial l_1}{\partial a} + m_3 \frac{\partial m_1}{\partial a} + n_3 \frac{\partial n_1}{\partial a} \right), \end{aligned} \right\} \quad (12)$$

Si l'on compare la seconde équation (12) avec l'expression (9) de la courbure de la courbe α dans la surface $\alpha\gamma$, on reconnaît que le second membre de (12) est le carré de la vitesse, multiplié par cette courbure, et par conséquent égal à la force centrifuge dans la direction opposée à $d\beta$, résultat qui, du reste, pouvait être prévu. Une remarque analogue peut être faite sur la troisième équation ci-dessus. Ces deux équations peuvent de plus au moyen de (11) être transformées en

$$\frac{\partial q}{\partial b} = \frac{h^2 \partial \varepsilon_1}{\varepsilon_1 \partial b}, \quad \frac{\partial q}{\partial c} = \frac{h^2 \partial \varepsilon_1}{\varepsilon_1 \partial c}. \quad (13)$$

L'équation de continuité (6) peut être transformée d'une manière analogue; mais on obtient plus vite le même résultat, si l'on remarque que, d'après la continuité, $h d\beta d\gamma$ doit être invariable pour tous les points d'une courbe décrite par un courant et par conséquent fonction de b et c seuls. Ainsi l'on peut poser

$$h = f(b, c) \varepsilon_2 \varepsilon_3. \quad (14)$$

Si maintenant nous considérons l'espèce particulière de mouvement pour laquelle est valable l'équation (3), $q = \frac{1}{2} h^2$, nous trouverons, en introduisant cette valeur de q dans les équations ci-dessus, que la première équation

tion devient identique et qu'inversement elle entraîne, par intégration, l'équation (1), tandis que les équations (13) donneront

$$\frac{1}{h} \frac{\partial h}{\partial b} = \frac{1}{\varepsilon} \frac{\partial \varepsilon_1}{\partial b}, \quad \frac{1}{h} \frac{\partial h}{\partial c} = \frac{1}{\varepsilon_1} \frac{\partial \varepsilon_1}{\partial c}, \quad (15)$$

d'où l'on déduit par intégration

$$h = F(a) \varepsilon_1, \quad (16)$$

F étant une fonction arbitraire.

Ce résultat qui, combiné avec (14), fournit les lois du mouvement et qui le détermine complètement, les conditions nécessaires étant données, peut être rendu plus intelligible par les considérations suivantes.

Le mouvement d'une particule d'eau peut être considéré comme une rotation autour d'un axe parallèle à $d\gamma$, passant par le centre de courbure du courant situé sur le prolongement de $d\beta$ et, de même, comme une rotation autour d'un axe parallèle à $d\beta$ et passant par l'autre centre de courbure, situé sur le prolongement de $d\gamma$. La vitesse angulaire de la première rotation est, si l'on se sert de la notation δ employée ci-dessus, $\frac{\delta}{dt}$ ou $h \frac{\delta}{da}$, si da est le chemin parcouru par une particule d'eau dans le temps dt . Si l'on s'imagine qu'une particule d'eau de la courbe a est liée par une ligne $d\beta$, normale à la courbe, avec la particule la plus proche de la surface $a\beta$, cette ligne se déplacera dans l'élément de temps dt et formera l'angle $-\frac{\delta h}{\partial \beta} dt$ avec sa position primitive, la direction positive étant comptée de la même manière que ci-dessus, pour la détermination de la courbure. La vitesse angulaire de la rotation de cette

ligne est donc

$$-\frac{\partial h}{\partial \beta} = -\frac{\partial h}{\partial b} \epsilon_2.$$

Mais, d'après l'équation (15), combinée avec (10), on a

$$-\frac{\partial h}{\partial b} \epsilon_2 = -h \frac{\epsilon_2}{\epsilon_1} \frac{\partial \epsilon_1}{\partial b} = -h \frac{\partial}{\partial a},$$

d'où ressort ce résultat remarquable, que cette vitesse angulaire est précisément égale et opposée à celle de la particule autour du centre de courbure situé sur le prolongement de $d\beta$. Une proposition analogue est valable pour la rotation autour de l'autre centre de courbure de la courbe du courant, et ces deux résultats, qui peuvent remplacer les équations (15), contiennent donc, conjointement avec l'équation de continuité, toutes les lois du mouvement du liquide. Un corps flottant librement dans le courant doit donc, en général, éprouver une rotation inverse de celle du courant, ce qui est précisément le contraire de ce qui arrive à un liquide, lorsqu'il éprouve une rotation uniforme autour d'un axe fixe.

Les expressions (14) et 16) de la vitesse peuvent être employées à la détermination graphique approximative, abstraction faite du frottement, d'un courant dans un canal. Si l'on suppose que le fond est plan, il peut être pris pour le plan coordonné fixe AB , et comme les courbes des courants ne diffèrent que peu du plan horizontal on peut, comme première approximation, supposer $d\gamma$ constant et égal à dc pour tous points de la surface AB . On a donc ici $\epsilon_s = 1$ et le mouvement sera d'après (14) et (15) déterminé par

$$h = F(a)\epsilon_1 = f(b)\epsilon_2.$$

Pour $b = 0$, on a $\varepsilon_1 = 1$, et pour $a = 0$, $\varepsilon_2 = 1$, de manière qu'on peut encore écrire ces équations sous la forme

$$h = [h]^{b=0} \frac{\partial a}{\partial \alpha} = [h]^{a=0} \frac{\partial b}{\partial \beta}.$$

Qu'on imagine des segments égaux à da et db portés sur les axes fixes, et inversement proportionnels aux vitesses, de manière qu'on ait

$$[h]^{b=0} da = [h]^{a=0} db;$$

alors on aura $da = d\beta$.

Les trajectoires orthogonales, dont l'une représente les directions des courants, formeront donc partout des carrés, si leurs traces sur les axes fixes sont marquées à intervalles infiniment petits, inversement proportionnels aux vitesses aux points considérés, et les côtés de ces carrés seront aussi inversement proportionnels aux vitesses.

Par conséquent, si les conditions limites nécessaires sont données, on pourra déterminer graphiquement la direction et la vitesse du courant; mais il va sans dire qu'on doit ici remplacer les carrés infiniment petits par des carrés petits, mais finis, et construire les trajectoires orthogonales de manière à satisfaire aux conditions données. Il faut pourtant remarquer que cette méthode de construction peut facilement entraîner une accumulation d'erreurs et par suite, conduire à des résultats incertains. C'est pourquoi l'on ne devra l'employer qu'avec précaution.

NOTES.

NOTE 1. On suppose qu'une partie de la surface est en repos et par conséquent plane, et qu'on fait passer le plan des xy par cette partie de la surface.

NOTE 2. La vitesse sera en chaque point dirigée suivant une droite rencontrant la verticale qui passe par le centre de l'orifice et ne dépendra que des distances du point considéré à cette verticale et à la surface libre.

NOTE 3. On ne désigne pas ici par $\frac{\partial a}{\partial a}$ la dérivée partielle par rapport à α , β et γ étant constants, mais la dérivée suivant la direction de $\alpha \cdot \frac{\partial b}{\partial \beta}$ et $\frac{\partial c}{\partial \gamma}$ ont des significations analogues.

NOTE 4. Lorenz parle ici et dans ce qui suit de la courbure d'une courbe dans une surface sur laquelle elle est tracée. Il résulte de son texte qu'il a en vue la courbure de la projection de la courbe sur le plan tangent à la surface au point considéré.

De même, quand il parle des divers centres de courbure d'une courbe, il entend les centres de courbure des projections de la courbe sur les différents plans qui passent par la tangente au point considéré.

NOTE 5. Lorenz dit qu'il prend pour courbes α les chemins parcourus par les particules du liquide.

C'est en général impossible; car on ne peut pas, en général, construire une surface qui passe par un point donné d'une des courbes et qui coupe ces courbes partout à angle droit: même dans le cas où l'on peut construire une famille de surfaces jouissant de ces propriétés, elles ne peuvent, en général, appartenir à un système de surfaces orthogonales, à moins que leur paramètre ne satisfasse à une certaine équation aux dérivées partielles.

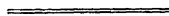
Si les surfaces de niveau $q = \text{const.}$ sont perpendiculaires aux courbes du courant, on doit avoir

$$\frac{\frac{\partial q}{\partial x}}{u} = \frac{\frac{\partial q}{\partial y}}{v} = \frac{\frac{\partial q}{\partial z}}{w},$$

ce qui entraîne

$$u\left(\frac{\partial v}{\partial z} - \frac{\partial w}{\partial y}\right) + v\left(\frac{\partial w}{\partial x} - \frac{\partial u}{\partial z}\right) + w\left(\frac{\partial u}{\partial y} - \frac{\partial v}{\partial x}\right) = 0,$$

équation qui est incompatible avec l'équation (5), à moins que les numérateurs de (5) ne s'évanouissent.



SUR LA RÉOLUTION

DES ÉQUATIONS ALGÈBRIQUES

AU MOYEN DE SÉRIES
ET D'INTÉGRALES DÉFINIES.

SUR LA RÉOLUTION DES ÉQUATIONS ALGÈBRIQUES AU MOYEN DE SÉRIES ET D'INTÉGRALES DÉFINIES.

TIDSSKRIFT FOR MATEMATIK 1867, P. 71—80.

Soit proposée l'équation

$$x = a + y\varphi(x), \quad (1)$$

à résoudre par rapport à x ; on pourra, comme on sait, se servir de la série de Lagrange

$$f(x) = f(a) + \varphi(a)f'(a)\frac{y}{1} + \dots + \frac{d^{n-1}(\varphi(a)^n f'(a))}{d a^{n-1}} \frac{y^n}{n!} + \dots \quad (2)$$

qui exprime non seulement x , mais une fonction quelconque $f(x)$ au moyen de a et de y . Dans l'application de cette série, on rencontrera pourtant une difficulté double, 1° parce que la série devient souvent divergente; 2° parce qu'elle ne donne qu'une seule racine de l'équation. C'est pourquoi il ne sera pas sans intérêt d'examiner la question plus en détail. Il me semble important avant tout d'obtenir des séries convergentes pour toutes les racines des équations algébriques, dont la résolution n'est en général qu'approximative.

Pour appliquer la série de Lagrange à la résolution de l'équation

$$x^a - x + a = 0, \quad (3)$$

seul type d'équation algébrique dont nous nous occupons ici, nous devons dans l'équation (2) poser

$$y = 1, \quad \varphi(a) = a^q,$$

et x , ou une puissance de x , sera donc déterminée par la série

$$x^p = a^p + p a^{p+q-1} + \dots + p \frac{(p+nq-1)(p+nq-2)\dots(p+n(q-1)+1)}{n!} a^{p+n(q-1)} + \dots \quad (4)$$

Si l'on se sert des signes de sommation et des fonctions Γ , cette série peut être écrite sous la forme

$$x^p = p \sum \frac{\Gamma(p+nq)}{\Gamma(p+n(q-1)+1) \Gamma(n+1)} a^{p+n(q-1)}, \quad (5)$$

n devant parcourir toutes les valeurs positives et entières de 0 jusqu'à ∞ . On peut en outre obtenir une expression analogue pour $\log x$ au moyen de l'équation ci-dessus, par exemple en faisant converger p vers 0. Si p est négatif, la formule est encore valable, pourvu que la fonction Γ soit définie de manière à vérifier encore la formule

$$\alpha \Gamma(\alpha) = \Gamma(\alpha+1)$$

pour α négatif; car alors les valeurs de la fonction pour $\alpha < 0$ sont déterminées par ses valeurs pour $\alpha > 0$.

La série trouvée de cette manière devient pourtant divergente aussitôt que la valeur numérique de a dépasse une certaine limite, et l'on trouve comme à l'ordinaire, que la condition de convergence est, abstraction faite du signe,

$$a^{q-1} < (q-1)^{q-1} : q^q. \quad (6)$$

De plus la série n'exprime que la racine qui s'évanouit en même temps que a , tandis que l'équation a encore $q-1$ racines différentes qui ne s'évanouissent pas avec a .

Pour trouver un développement convergent qui représente toutes les racines de l'équation, nous partirons de l'équation plus générale

$$1 = \frac{y}{\varphi(x)} + \frac{z}{\phi(x)}, \quad (7)$$

d'où nous déduirons pour toute fonction de x deux séries, procédant l'une suivant les puissances croissantes de y , l'autre suivant celles de z . Ces séries peuvent être obtenues au moyen de la série de Lagrange, généralisée par Laplace; aussi peut-on dire que nos résultats sont contenus dans cette formule, bien qu'on n'ait pas encore, que je sache, fait cette remarque.

Cependant, comme la forme symétrique de l'équation ci-dessus rend les calculs considérablement plus faciles, je préfère donner une démonstration directe des séries cherchées.

En différentiant l'équation donnée par rapport à y et z , on trouve

$$\begin{aligned} 0 &= \frac{1}{\varphi(x)} - \left(\frac{y\varphi'(x)}{\varphi(x)^2} + \frac{z\phi'(x)}{\phi(x)^2} \right) \frac{\partial x}{\partial y}, \\ 0 &= \frac{1}{\phi(x)} - \left(\frac{y\varphi'(x)}{\varphi(x)^2} + \frac{z\phi'(x)}{\phi(x)^2} \right) \frac{\partial x}{\partial z}, \end{aligned}$$

et par conséquent

$$\varphi(x) \frac{\partial x}{\partial y} = \phi(x) \frac{\partial x}{\partial z}. \quad (8)$$

Si nous avons de plus deux fonctions $f(x)$ et $F(x)$ développables en série suivant les puissances croissantes de z , savoir

$$\begin{aligned} f(x) &= A_0 + A_1 \frac{z}{1} + A_2 \frac{z^2}{1 \cdot 2} + \dots, \\ F(x) &= B_0 + B_1 \frac{z}{1} + B_2 \frac{z^2}{1 \cdot 2} + \dots, \end{aligned}$$

on obtiendra en différentiant respectivement par rapport à z et y

$$f'(x) \frac{\partial x}{\partial z} = A_1 + A_2 z + \dots,$$

$$F'(x) \frac{\partial x}{\partial y} = \frac{dB_0}{dy} + \frac{dB_1}{dy} \frac{z}{1} + \dots$$

Si la fonction F est déterminée de manière à rendre ces deux séries identiques, on aura

$$A_1 = \frac{dB_0}{dy}, \quad A_2 = \frac{dB_1}{dy}, \quad \dots,$$

et

$$f'(x) \frac{\partial x}{\partial z} = F'(x) \frac{\partial x}{\partial y},$$

d'où l'on tire, par comparaison avec l'équation (8),

$$f'(x) \varphi(x) = F'(x) \phi(x). \quad (9)$$

Les coefficients A_0 et B_0 peuvent être exprimés par

$$A_0 = f(\beta), \quad B_0 = F(\beta),$$

où β désigne la valeur de x pour $z = 0$. β est donc d'après (7) une racine de l'équation

$$\varphi(\beta) = y. \quad (10)$$

Nous avons donc

$$A_1 = \frac{dB_0}{dy} = F'(\beta) \frac{d\beta}{dy} = \frac{\varphi(\beta)}{\phi(\beta)} \frac{df(\beta)}{dy},$$

et nous pouvons déterminer B_1 , car ce coefficient dépend de la même manière de F que A_1 de f ; c'est pourquoi nous aurons

$$B_1 = \frac{\varphi(\beta)}{\phi(\beta)} \frac{dF(\beta)}{dy} = \left(\frac{\varphi(\beta)}{\phi(\beta)} \right)^2 \frac{df(\beta)}{dy}.$$

Ensuite la valeur de A_2 se déduira de celle de B_1 , car A_2 est la dérivée de B_1 par rapport à y . En continuant de cette manière, on trouvera facilement les termes successifs de la série $f(x)$ et l'on obtiendra

$$f(x) = f(\beta) + \frac{\varphi(\beta)}{\phi(\beta)} \frac{df(\beta)}{dy} \cdot \frac{z}{1} \cdots \frac{d^{n-1}}{dy^{n-1}} \left[\left(\frac{\varphi(\beta)}{\phi(\beta)} \right)^n \frac{df(\beta)}{dy} \right] \frac{z^n}{n!} \cdots \quad (11)$$

Comme on peut, dans l'équation proposée, permuter en même temps y et z , φ et ϕ , on déduira de la dernière équation, si γ est une racine de l'équation

$$\varphi(\gamma) = z, \quad (10')$$

en effectuant la permutation indiquée,

$$f(x) = f(\gamma) + \frac{\phi(\gamma)}{\varphi(\gamma)} \frac{df(\gamma)}{dz} \frac{y}{1} + \cdots + \frac{d^{n-1}}{dz^{n-1}} \left[\left(\frac{\phi(\gamma)}{\varphi(\gamma)} \right)^n \frac{df(\gamma)}{dz} \right] \frac{y^n}{n!} \cdots \quad (11')$$

Si nous considérons l'équation algébrique (3), nous reconnaitrons qu'elle peut, de trois manières différentes, être rapportée au type (7), car on peut isoler chacun de ses trois termes à gauche du signe d'égalité, puis diviser l'équation par ce membre.

On reconnaît pourtant que c'est seulement en donnant à l'équation la forme

$$1 = \frac{1}{x^{q-1}} - \frac{a}{x^q},$$

que nous obtiendrons toutes les solutions, à l'exception d'une seule, donnée par l'équation (5).

Nous devons donc poser

$$\varphi(x) = x^{q-1}, \quad \phi(x) = x^q, \quad y = 1, \quad z = -a,$$

et d'après (10) et (10')

$$\beta = y^{\frac{1}{q-1}} \quad \text{et} \quad \gamma = z^{\frac{1}{q}}.$$

Les deux séries de x^p deviendront donc d'après (11) et (11')

$$x^p = y^{\frac{p}{q-1}} + y^{\frac{-1}{q-1}} \frac{d\left(y^{\frac{p}{q-1}}\right)}{dy} \frac{(-a)}{1} + \dots + \frac{d^{n-1}}{dy^{n-1}} \left[y^{-\frac{n}{q-1}} \frac{d\left(y^{\frac{p}{q-1}}\right)}{dy} \right] \frac{(-a)^n}{n!} + \dots,$$

où l'on doit poser $y = 1$ après avoir exécuté la différentiation, et

$$x^p = z^{\frac{p}{q}} + z^{\frac{1}{q}} \frac{d\left(z^{\frac{p}{q}}\right)}{dz} \frac{1}{1} + \dots + \frac{d^{n-1}}{dz^{n-1}} \left[z^{\frac{n}{q}} \frac{d\left(z^{\frac{p}{q}}\right)}{dz} \right] \frac{1}{n!} + \dots,$$

où l'on pose $z = -a$.

Si l'on conserve la notation β pour la racine $(q-1)^{\text{ième}}$ de l'unité et si l'on pose pour abrégé

$$\frac{nq-p}{q-1} = r_n, \quad \frac{n+p}{q} = s_n,$$

la première série donnera

$$\left(\frac{x}{\beta}\right)^p = 1 + r_0 \frac{a}{\beta} + \dots + r_0 \frac{(r_n-1)(r_n-2) \dots (r_n-n+1)}{1 \cdot 2 \dots n} \left(\frac{a}{\beta}\right)^n + \dots, \quad (13)$$

tandis que l'autre donnera

$$x^p = (-a)^{s_0} + s_0(-a)^{s_1-1} + \dots + s_0 \frac{(s_n-1)(s_n-2) \dots (s_n-n+1)}{1 \cdot 2 \dots n} (-a)^{s_n-n} \dots + (14)$$

Pour obtenir les conditions de convergence de ces deux séries, le mieux est de comparer le $n^{\text{ième}}$ terme avec le $(n+q-1)^{\text{ième}}$ de la première et le $(n+q)^{\text{ième}}$ de la seconde: de cette manière, on trouvera, comme à l'ordinaire, que la série (13) est convergente sous les mêmes conditions que (5), c'est-à-dire si l'on a, abstraction faite du signe,

$$a < (q-1)q^{\frac{q}{1-q}},$$

tandis que la série (14) est précisément convergente quand cesse la convergence des autres séries, savoir quand

$$a > (q-1)q^{\frac{q}{1-q}}.$$

Les séries (5) et (13) se suppléent, car la dernière donne les $(q-1)$ valeurs de x^p qui ne sont pas contenues dans (5), si l'on fait prendre à β les diverses valeurs de la racine $(q-1)^{\text{ième}}$ de l'unité. Au contraire, aussitôt que a dépasse la limite indiquée ci-dessus, on ne peut se servir que de la série (14) et elle donne encore la solution complète, car $(-a)$ est élevé à des puissances dont les exposants sont des nombres fractionnaires de dénominateur q : on aura donc q solutions différentes avec les q valeurs que prend $(-a)^{\frac{1}{q}}$.

Si a est un nombre réel, on reconnaît, si l'on peut se servir des premières séries, a ne dépassant pas la limite indiquée ci-dessus, que l'équation aura deux ou trois racines réelles selon que son degré est pair ou impair, β ayant toujours une ou deux valeurs réelles et la série (5) donnant toujours une valeur réelle. Si, au contraire, a est positif et dépasse la limite, on n'aura d'après (14) aucune racine réelle, si le degré de l'équation est pair; on en aura une si ce degré est impair. Cette disparition de deux racines réelles fait soupçonner que l'équation a des racines égales pour la valeur limite de a ,

$$a = (q-1)q^{\frac{q}{1-q}},$$

ce qui peut servir à la résolution de l'équation dans ce cas limite. Or on doit avoir, dans le cas des racines égales,

$$qx^{q-1} - 1 = a \text{ et par suite } x = q^{\frac{1}{1-q}},$$

valeur qui est précisément racine de l'équation proposée, quand a prend la valeur en question.

Il n'est pas sans intérêt de remarquer qu'on peut dans la série (13) réunir dans une somme tous les termes qui contiennent comme facteur commun la même puissance de $\beta = \frac{1}{q-1}$, car ce facteur se répète pour chaque $(q-1)^{\text{ième}}$ terme, et chaque somme prise séparément peut être considérée comme la racine $(q-1)^{\text{ième}}$ d'un nombre réel. Dans la série (14), les termes qui correspondent à une valeur entière de s_n s'évanouiront, et, si elle est ordonnée d'une manière analogue, cette série contiendra $(q-1)$ sommes dont chacune, prise séparément, peut être considérée comme la racine $q^{\text{ième}}$ d'un nombre réel.

Si l'on veut exprimer la série (14) au moyen des fonctions Γ , on peut le faire de la manière suivante. Le terme général du développement contient le produit

$$(s_n-1)(s_n-2) \dots (s_n-m) \dots (s_n-n+1),$$

dont le premier facteur est toujours positif, le dernier négatif pour des valeurs suffisamment grandes de n , de manière qu'on peut toujours trouver un nombre entier m pour lequel s_n-m est compris entre 0 et 1. Les $n-m-1$ facteurs suivants, qui par conséquent sont négatifs, peuvent, si l'on change leurs signes, être comme les précédents exprimés par les fonctions Γ , et le produit peut s'écrire sous la forme

$$\frac{\Gamma(s_n) \Gamma(n-s_n)}{\Gamma(s_n-m) \Gamma(1-s_n+m)} (-1)^{n-m-1}.$$

Le dénominateur peut être transformé au moyen de la formule

$$\Gamma(n) \Gamma(1-n) = \frac{\pi}{\sin n\pi},$$

et, comme on le voit facilement, la série (14) prendra la forme

$$x^p = \sum_{s_0} \frac{\sin s_n \pi}{\pi} \frac{\Gamma(s_n) \Gamma(n-s_n)}{\Gamma(n+1)} (-1)^{s_n-1} a^{s_n-n}. \quad (15)$$

Dans cette formule m n'entre pas et elle est valable pour toutes les valeurs que prend n , de 0 jusqu'à ∞ , si l'on a égard à la remarque faite ci-dessus relativement à la définition de la fonction Γ pour un argument négatif.

Si l'on pose

$$n = mq + \mu,$$

m et μ étant des nombres entiers, on aura $s_n = m + s_\mu$ et l'on obtiendra la somme de tous les termes en faisant prendre à μ toutes les valeurs de 0 à $q-1$ et à m toutes les valeurs de 0 à ∞ . Si l'on différencie par rapport à a , la série prendra la forme

$$\left. \begin{aligned} \frac{d}{da} x^p &= \sum_{\mu=0}^{\mu=q-1} s_0 \frac{\sin s_\mu \pi}{\pi} (-1)^{s_\mu} a^{s_\mu-\mu-1} R_\mu, \\ R_\mu &= \sum \frac{\Gamma(m+s_\mu) \Gamma(m(q-1)-s_\mu+\mu+1)}{\Gamma(mq+\mu+1)} a^{-m(q-1)} \end{aligned} \right\} \quad (16)$$

La première somme ne contient que $(q-1)$ termes; car, d'après la remarque faite ci-dessus, le terme pour lequel s_μ est un nombre entier s'évanouit. On a ici cherché une expression de la dérivée de x^p par rapport à a , parce que cette série peut, comme nous le verrons, s'exprimer plus facilement par des intégrales définies. Du reste, on voit facilement qu'on n'a pas besoin d'intégrer pour obtenir x^p : car, en différenciant l'équation proposée, on trouve

$$(qx^{q-1} - 1) \frac{dx}{da} + 1 = 0,$$

équation qui, combinée avec la proposée, donne

$$\frac{dx}{da} = \frac{x}{aq - x(q-1)},$$

d'où l'on peut tirer x au moyen de $\frac{dx}{da}$.

Les séries trouvées ici peuvent toutes être facilement sommées au moyen d'intégrales définies.

Pour sommer la série R_μ dans la formule (16), on peut, par exemple, se servir de l'intégrale eulérienne

$$\int_0^1 u^{a-1} (1-u)^{\beta-1} du = \frac{\Gamma(a) \Gamma(\beta)}{\Gamma(a+\beta)},$$

par où l'on obtient immédiatement

$$R_\mu = \sum \int_0^1 u^{m+s\mu-1} (1-u)^{m(q-1)-s\mu+\mu} du \cdot a^{-m(q-1)},$$

ou, en effectuant la sommation,

$$R_\mu = a^{q-1} \int_0^1 \frac{u^{s\mu-1} (1-u)^{\mu-s\mu}}{a^{q-1} - u(1-u)^{q-1}} du.$$

Si l'on pose ici

$$u = \frac{z-a}{z},$$

on obtiendra

$$R_\mu = a^{\mu+1-s\mu} \int_a^\infty \frac{z^{q-\mu-1} (z-a)^{s\mu-1}}{z^q - z + a} dz,$$

où z croît numériquement de telle sorte que la limite supérieure de l'intégrale est $+\infty$ si a est positif et $-\infty$ si a est négatif. Par cette transformation on obtiendra

l'intégrale d'une fraction dont la décomposition exige la solution de l'équation

$$z^q - z + a = 0,$$

ce qui est l'équation proposée. Autant que je sache, les différentes solutions qu'on a trouvées jusqu'ici des équations algébriques au moyen d'intégrales définies peuvent être rapportées à ce type. Mais une telle solution ne peut pas être considérée comme une solution effective, car on ne pourrait pas, par exemple, trouver une solution d'une équation du troisième degré, et ces intégrales définies ne sont d'aucune utilité tant qu'on n'en peut pas déduire un développement en série.

Un type nouveau de solutions peut être obtenu par la sommation de la série (5) au moyen d'une intégrale définie. J'emploierai à cet effet le théorème suivant de Gauss et Legendre:

$$\Gamma(am) = \Gamma(a) \Gamma\left(a + \frac{1}{m}\right) \dots \Gamma\left(a + \frac{m-1}{m}\right) m^{am-\frac{1}{2}} (2\pi)^{\frac{1-m}{2}}.$$

et je développerai l'application relative aux équations du troisième et du quatrième degré.

Pour $q = 3$, $p = 1$, la formule (5) donnera

$$x = \sum \frac{\Gamma(3n+1)}{\Gamma(2n+2) \Gamma(n+1)} a^{2n+1},$$

où l'on pose, d'après le théorème cité,

$$\Gamma(3n+1) = \frac{1}{2\pi} \Gamma\left(n + \frac{1}{3}\right) \Gamma\left(n + \frac{2}{3}\right) \Gamma(n+1) 3^{3n+\frac{1}{2}}.$$

Après avoir différentié par rapport à a , on aura

$$\frac{dx}{da} = \frac{1}{2\pi} \sum \frac{\Gamma\left(n + \frac{1}{3}\right) \Gamma\left(n + \frac{2}{3}\right)}{\Gamma(2n+1)} 3^{3n+\frac{1}{2}} a^{2n},$$

et la sommation peut être effectuée comme ci-dessus au moyen de l'intégrale eulérienne.

De cette manière on obtient

$$\frac{dx}{da} = \frac{\sqrt{3}}{2\pi} \int_0^1 \frac{u^{-\frac{2}{3}}(1-u)^{-\frac{1}{3}}}{1-27a^2u(1-u)} du,$$

qui se transformera en

$$\frac{dx}{da} = \frac{\sqrt{3}}{2\pi} \int_0^\infty \frac{z^{-\frac{2}{3}}(z+1)}{(1+z)^2-27a^2z} dz$$

par la substitution

$$u = \frac{z}{1+z}.$$

On pourra sans difficulté effectuer les intégrations, en décomposant la fraction et en appliquant la formule

$$\int_0^\infty \frac{u^{a-1}}{1+u} du = \frac{\pi}{\sin a\pi}, \quad 0 < a < 1,$$

par où l'on trouvera les expressions algébriques bien connues des racines de l'équation. L'intégrale donnera la solution complète, bien que la série, dont la sommation a fourni l'intégrale considérée, ne corresponde qu'à une seule racine de l'équation.

Pour l'équation du quatrième degré, la formule (5) donnera

$$x = \sum \frac{\Gamma(4n+1)}{\Gamma(3n+2)\Gamma(n+1)} a^{2n+1},$$

et le théorème de Gauss permettra d'écrire

$$\Gamma(4n+1) = \Gamma(2n+\frac{1}{2}) \Gamma(2n+1) 2^{4n+\frac{1}{2}} (2\pi)^{-\frac{1}{2}}$$

et

$$\Gamma(2n+1) = \Gamma(n+\frac{1}{2}) \Gamma(n+1) 2^{2n+\frac{1}{2}} (2\pi)^{-\frac{1}{2}}.$$

En multipliant ces deux équations, on obtiendra une expression de $\Gamma(4n+1)$ qu'on introduira dans l'expression ci-dessus de x ; puis, après avoir différentié l'équation par rapport à a , on aura

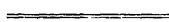
$$\frac{dx}{da} = \frac{1}{\pi} \sum \frac{\Gamma(2n + \frac{1}{2}) \Gamma(n + \frac{1}{2})}{\Gamma(3n + 1)} 2^{6n} a^{3n}.$$

Comme ci-dessus, la sommation peut être exécutée au moyen d'une intégrale définie et l'on obtiendra

$$\frac{dx}{da} = \frac{1}{\pi} \int_0^1 \frac{u^{-\frac{1}{2}} (1-u)^{-\frac{1}{2}}}{1 - 64a^3 u^2 (1-u)} du,$$

où l'intégration dépend de la résolution d'une équation du troisième degré et peut être exécutée comme ci-dessus.

Tandis qu'on peut, de cette manière, résoudre effectivement les équations du troisième et du quatrième degré au moyen d'intégrales définies, le procédé échoue, au contraire, pour les équations de degré supérieur, à raison d'intégrations qui ne peuvent pas être exécutées, ainsi qu'on pouvait du reste s'y attendre.



CONTRIBUTION
A LA THÉORIE DES NOMBRES.

CONTRIBUTION A LA THÉORIE DES NOMBRES.

TIDSSKRIFT FOR MATEMATIK, 1871, P. 97—114.

Si l'on cherche à résoudre le problème de trouver le nombre des solutions de l'équation indéterminée

$$m^2 + cn^2 = N, \quad (1)$$

où m et n désignent des nombres entiers positifs, nuls ou négatifs qui satisfont à l'équation, c et N étant des entiers positifs, il arrive qu'on peut parfois trouver, pour certaines valeurs singulières de c , une loi simple mettant en évidence la dépendance qui rattache le nombre en question aux facteurs de N . Ici et dans ce qui suit, par ces mots „le nombre des solutions“, nous entendons, à moins qu'une autre signification ne leur soit expressément donnée, le nombre total des valeurs, positives négatives et nulles, que peuvent prendre m et n dans l'équation (1).

Si nous considérons, par exemple, l'équation

$$m^2 + n^2 = N,$$

et si nous posons successivement

$$m = 0,$$

$$n = \pm 1, \pm 2, \dots,$$

$$m = \pm 1,$$

$$n = 0, \pm 1, \pm 2, \dots,$$

lées différentes valeurs que N prendra dans ces conditions seront énumérées dans le tableau suivant

$$1, 2, 4, 5_2, 8, 9, 10_2, 13_2, 16, 17_2, 18, 20_2, \dots,$$

où les nombres qui n'ont pas d'indice correspondent aux valeurs de N , pour lesquelles le nombre de solutions est 4; ce nombre étant le double, c'est-à-dire égal à 8, pour les nombres qui ont l'indice 2.

On peut, à titre d'essai, chercher à former la même série de la manière suivante. On commence par écrire la série des nombres entiers :

$$1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17, 18, \dots$$

Puis on soustrait 3, qui est le premier nombre qui ne se trouve pas dans la série ci-dessus, et tous les multiples de 3, par conséquent

$$-3, -6, -9, -12, -15, -18, \dots,$$

nombres qui sont affectés du signe négatif pour indiquer qu'ils doivent être soustraits.

Puis nous formons les séries

$$\begin{array}{ccccccc} 5, & 10, & 15, & 20, & \dots, \\ -7, & -14, & & \dots \end{array}$$

et ainsi de suite, de manière que chaque série successive commence par le nombre impair suivant, changé de signe, et soit continuée par les multiples du premier nombre.

En additionnant les nombres correspondants de ces séries, on obtiendra la série

$$1, 2, 4, 5_2, 8, 9, 10_2, 13_2, 16, 17_2, 18, 20_2, \dots$$

où l'indice 2 signifie que le nombre correspondant se

trouve deux fois dans les séries en question. Nous retrouverons ainsi, au moins pour les nombres considérés, la première série; et, en continuant de la même manière, on obtiendrait facilement une plus grande vraisemblance de la concordance absolue des deux séries. Si cette concordance existe réellement, ce qu'on ne peut démontrer en toute rigueur par cette méthode empirique, nous pourrions facilement trouver la loi du nombre des solutions, parce que les dernières séries sont formées d'après une loi simple. Le terme général de la première série peut être désigné par n , celui de la troisième par $5n$, celui de la cinquième par $9n$ etc., par conséquent en général par $(4m+1)n$, tandis que le terme général des séries à termes négatifs devient $(4m+3)n$, où m peut prendre toutes les valeurs entières à partir de zéro jusqu'à l'infini, et n de 1 jusqu'à l'infini. Un nombre N se trouvera par conséquent autant de fois dans la série finale qu'il y a d'unités dans le nombre des diviseurs de N de la forme $4m+1$, diminué du nombre de ceux qui ont la forme $4m+3$. Soit ρ_N le nombre des solutions de l'équation $m^2+n^2=N$, et soient a_N et b_N les nombres des diviseurs de N de la forme $4m+1$ et $4m+3$; ρ_N est déterminé par

$$\rho_N = 4(a_N - b_N).$$

On peut, en fait, par ce procédé complètement empirique, obtenir une solution exacte de la plupart des problèmes qui seront résolus dans ce qui suit; mais il va sans dire que de telles solutions ne sont pas satisfaisantes au point de vue mathématique, parce qu'elles se présentent comme tout à fait fortuites et manquent de la rigueur qu'on est accoutumé à demander aux

résultats mathématiques. Cette méthode d'essai n'est pourtant pas sans portée; c'est pourquoi je ne l'ai pas négligée, parce qu'elle peut, par l'orientation des résultats qu'on cherche, indiquer la marche qu'il faut suivre dans les calculs.

Si l'on cherche à démontrer toutes les propositions qu'on peut trouver par la méthode d'essai indiquée ci-dessus, on pourra, ou bien suivre la méthode plus synthétique dont on se sert d'ordinaire dans la théorie des nombres, ou bien chercher à obtenir la solution par une voie purement analytique; c'est cette dernière que je préférerai ici.* La méthode analytique consiste dans la transformation de la série

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+cn^2}, \quad (q < 1)$$

en une série d'une autre forme. Puis, en comparant les termes des deux séries qui ont le même exposant N , on détermine le nombre de fois qu'on rencontre cet exposant et par là le nombre des solutions de l'équation $m^2 + cn^2 = N$.

C'est de cette manière que Jacobi (Journal de Crelle, tome 12) a résolu le problème dans le cas de $c = 1$, en transformant la série au moyen de la théorie des fonctions elliptiques, Dans sa „Théorie des nombres“ (3^{ième} éd., p. 216), Legendre, en établissant que tout nombre positif est une somme de quatre carrés (ou

* Dans le dernier fasc. de ce journal, que j'ai reçu après avoir terminé la rédaction de ce mémoire, M. J. Petersen a traité le même problème par la première méthode. On aura donc l'occasion de comparer les deux méthodes.

moins) appelle l'attention sur ce fait que ce théorème peut être démontré au moyen des séries de la théorie des fonctions elliptiques, et il ajoute que la démonstration de l'identité des séries, et par suite du théorème, doit pouvoir être faite au moyen de calculs purement analytiques, par où l'on obtiendrait „la démonstration la plus simple qu'il soit possible“ de donner du théorème en question.

C'est dans cette conviction que j'ai essayé de faire ressortir *une solution immédiate, analytique et plus générale*, de cette sorte de problèmes, de manière à éviter le détour difficile des fonctions elliptiques, qui, de plus, ne peuvent être employées que dans le cas singulier mentionné plus haut. De plus, j'indiquerai une application des résultats trouvés à *la sommation de quelques séries*, problème qui a été l'origine de ces recherches.

Nous considérerons en premier lieu la série

$$\sum_{-\infty}^{+\infty} q^{n^2} x^n = f(x), \quad (2)$$

où n peut prendre toutes les valeurs positives, négatives et nulles, et, après avoir déterminé l'équation fondamentale de la série considérée comme fonction de x , nous transformerons la série au moyen de cette équation en un produit d'un nombre infiniment grand de facteurs. La solution de ce problème est, sous une forme peu différente, une question bien connue dans la théorie des fonctions elliptiques; mais comme je ne présuppose ici aucune connaissance de cette théorie, j'indiquerai en détail le calcul.

En supposant que q est plus petit que l'unité, la série est convergente pour une valeur quelconque de x .

Si l'on remplace x par q^2x , on obtiendra

$$f(q^2x) = \sum_{-\infty}^{+\infty} q^{n^2+2n}x^n = \sum_{-\infty}^{+\infty} q^{(n+1)^2-1}x^n,$$

expression dans laquelle $n+1$ peut être remplacé par n . Par conséquent, on aura

$$f(q^2x) = \frac{1}{qx} f(x). \quad (3)$$

Inversement, si l'on cherchait à déterminer $f(x)$ au moyen de cette équation, on pourrait le faire en exprimant cette fonction par une série procédant suivant les puissances positives et négatives de x et en déterminant les coefficients au moyen de l'équation fondamentale. En procédant de cette manière, on se convaincra que tous les coefficients peuvent être déterminés, à un facteur constant près, qui reste arbitraire. Toute fonction $F(x)$ qui satisfait à l'équation (3) peut donc être exprimée d'une manière générale par

$$\varphi(q) F(x) = f(x), \quad (4)$$

où φ est une fonction arbitraire, dépendant seulement de q .

Or, la fonction

$$F(x) = (1+qx)(1+q^3x) \dots (1+qx^{-1})(1+q^3x^{-1}) \dots$$

satisfait à l'équation fondamentale; car, en remplaçant x par q^2x , on obtiendra

$$F(q^2x) = (1+q^3x)(1+q^5x) \dots \\ (1+q^{-1}x^{-1})(1+qx^{-1})(1+q^3x^{-1}) \dots,$$

d'où il s'ensuit que

$$F(q^2x) = \frac{1+q^{-1}x^{-1}}{1+qx} F(x) = \frac{1}{qx} F(x).$$

Nous aurons, par suite, en nous servant d'une notation abrégée,

$$\sum_{-\infty}^{+\infty} q^{n^2} x^n = \varphi(q) \prod_1^{\infty} (1+q^{2h-1}x)(1+q^{2h-1}x^{-1}), \quad (5)$$

où $\prod_1^{\infty}$ désigne le produit d'une suite infinie de facteurs obtenus en posant successivement $h = 1, 2, \dots$ jusqu'à $h = \infty$. A présent, la fonction indéterminée $\varphi(q)$ doit être déterminée par comparaison des valeurs que prennent les deux membres de l'équation (4), quand on y remplace x par des valeurs données. Pour $x = -1$ et $x = \sqrt{-1}$, on obtiendra les deux relations

$$1 - 2q + 2q^4 - 2q^9 \dots = \varphi(q) \prod_1^{\infty} (1 - q^{2h-1})^2,$$

$$1 - 2q^4 + 2q^{16} - 2q^{36} \dots = \varphi(q) \prod_1^{\infty} (1 + q^{4h-2}),$$

où la seconde série est déduite de la première par le changement de q en q^4 . Nous obtiendrons donc

$$\varphi(q^4) \prod_1^{\infty} (1 - q^{8h-4})^2 = \varphi(q) \prod_1^{\infty} (1 + q^{4h-2}),$$

d'où résulte

$$\frac{\varphi(q)}{\varphi(q^4)} = \prod_1^{\infty} \frac{(1 - q^{8h-4})^2}{1 + q^{4h-2}} = \prod_1^{\infty} (1 - q^{4h-2})(1 - q^{8h-4}).$$

Dans la dernière suite de facteurs, q aura successivement tous les exposants 2, 6, 10, ..., 4, 12, 20, ... c'est-à-dire tous les nombres pairs, à l'exception de 8, 16, 24, ...; pour cette raison l'équation peut encore s'écrire

$$\frac{\varphi(q)}{\varphi(q^4)} = \prod_1^{\infty} \frac{1 - q^{2h}}{1 - q^{8h}}.$$

Dans cette équation, on remplace q par q^4 , et l'équation

$$\frac{\varphi(q^4)}{\varphi(q^{16})} = \prod_1^{\infty} \frac{1 - q^{8h}}{1 - q^{32h}},$$

formée de cette manière, donnera, par multiplication avec la précédente,

$$\frac{\varphi(q)}{\varphi(q^{16})} = \prod_1^{\infty} \frac{1 - q^{2h}}{1 - q^{32h}}.$$

Ici on pourra de nouveau remplacer q par q^{16} et ainsi de suite, et, en remarquant que $q < 1$ et qu'on aura $q^{\infty} = 0$ et, d'après (5), $\varphi(0) = 1$, on obtiendra finalement

$$\varphi(q) = \prod_1^{\infty} (1 - q^{2h}),$$

par où la fonction φ est déterminée.

Si l'on introduit cette expression dans l'équation (5), et si l'on pose en même temps $x = 1$, le seul cas dont nous nous servirons ici, on aura

$$\sum_{-\infty}^{+\infty} q^{n^2} = \prod_1^{\infty} (1 - q^{2h})(1 + q^{2h-1})^2.$$

Cette équation acquiert une forme plus convenable pour notre but si on la multiplie par le facteur

$$\begin{aligned} 1 &= \prod_1^{\infty} \frac{(1 - q^{4h-2})(1 - q^{4h})}{(1 - q^{4h-2})(1 - q^{4h})} \\ &= \prod_1^{\infty} \frac{(1 - q^{2h})}{(1 - q^{4h-2})(1 - q^{4h})} = \prod_1^{\infty} \frac{1}{(1 - q^{4h-2})(1 + q^{2h})}, \end{aligned}$$

par où l'on obtient finalement

$$\sum_{-\infty}^{+\infty} q^{n^2} = \prod_1^{\infty} \frac{(1-q^{2h})(1+q^{2h-1})}{(1+q^{2h})(1-q^{2h-1})}. \quad (6)$$

A présent il est facile de former le produit infini qui exprime la somme double

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+cn^2} = \sum_{-\infty}^{+\infty} q^{m^2} \sum_{-\infty}^{+\infty} q^{cn^2},$$

savoir

$$\begin{aligned} & \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+cn^2} \\ &= \prod_1^{\infty} \frac{(1-q^{2h})(1+q^{2h-1})}{(1+q^{2h})(1-q^{2h-1})} \frac{(1-q^{2hc})(1+q^{(2h-1)c})}{(1+q^{2hc})(1-q^{(2h-1)c})}. \end{aligned} \quad (7)$$

Il reste encore à résoudre le problème de transformer le produit infini en une série de puissances, mais toutefois différente de la série primitive.

J'ai cherché à obtenir cette transformation au moyen de certains produits infinis de facteurs fractionnaires, en décomposant ces fractions d'après les règles ordinaires de la décomposition des fractions rationnelles. Nous considérerons en premier lieu le produit

$$\frac{1}{2-x} \prod_1^{\infty} \frac{1-q^{2h-1}x+q^{4h-2}}{1-q^{2h}x+q^{4h}} = P(x). \quad (8)$$

Cette fraction donne par sa décomposition une série de la forme suivante

$$P(x) = \frac{A_0}{2-x} + \frac{A_1}{1-q^2x+q^4} + \frac{A_2}{1-q^4x+q^8} + \dots,$$

où l'on déterminera A_0 en faisant $x=2$ dans $P(x)(2-x)$, puis A_1 en faisant $x = \frac{1+q^4}{q^2}$ dans $(1-q^2x+q^4)P(x)$,

et ainsi de suite. On trouvera de cette manière

$$A_0 = \prod_1^{\infty} \left(\frac{1 - q^{2^{h-1}}}{1 - q^{2^h}} \right)^2,$$

$$A_1 = A_0 q(1 + q^2), \quad A_2 = A_0 q^2(1 + q^4), \quad A_3 = A_0 q^3(1 + q^6), \quad \dots$$

Par conséquent, on aura

$$\frac{P(x)}{A_0} = \frac{1}{2-x} + \frac{q(1+q^2)}{1-q^2x+q^4} + \frac{q^2(1+q^4)}{1-q^4x+q^8} + \dots \quad (9)$$

Si l'on pose $x = -2$, on obtiendra

$$\frac{4P(-2)}{A_0} = \prod_1^{\infty} \left(\frac{1 + q^{2^{h-1}}}{1 + q^{2^h}} \cdot \frac{1 - q^{2^h}}{1 - q^{2^{h-1}}} \right)^2,$$

équation dont le second membre est identique avec le second membre de l'équation (7), où l'on fait $c = 1$. Pour ces valeurs de c et x , les équations (7) et (9) donneront

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+n^2} = 1 + 4 \left(\frac{q}{1+q^2} + \frac{q^2}{1+q^4} + \dots \right). \quad (10)$$

Si l'on se sert du signe de sommation, la dernière équation peut être écrite

$$\begin{aligned} \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+n^2} &= 1 + 4 \sum_1^{\infty} \frac{q^n}{1+q^{2n}} \\ &= 1 + 4 \sum_1^{\infty} \left[q^n - q^{3n} + q^{5n} \dots \right], \end{aligned}$$

ou, si l'on se sert d'un nouveau signe de sommation double,

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+n^2} \\ = 1 + 4 \sum_0^{\infty} \sum_1^{\infty} \left[q^{(4m+1)n} - q^{(4m+3)n} \right]; \quad (11)$$

dans le second membre m parcourt toutes les valeurs à partir de 0 jusqu'à ∞ , et n à partir de 1 jusqu'à ∞ .

Nous obtiendrons ainsi le résultat mentionné ci-dessus, savoir, que si l'équation

$$m^2 + n^2 = N \quad (12)$$

a ρ_N solutions, et si les équations

$$(4m+1)n = N \quad \text{et} \quad (4m+3)n = N$$

ont respectivement a_N et b_N solutions, n et m n'ayant que des valeurs positives (et de plus m la valeur zéro) dans les dernières équations, on aura d'après (11)

$$\rho_N = 4(a_N - b_N),$$

puisque q doit se trouver le même nombre de fois élevé à la même puissance N dans les deux membres de l'équation (11). On reconnaît facilement que a_N et b_N sont les nombres de diviseurs de N des formes $4m+1$ et $4m+3$. Ce résultat peut être énoncé par la proposition suivante:

Si le nombre N contient les facteurs premiers $p_1, p_2, \dots$ de la forme $4m+1$ avec les exposants $a_1, a_2, \dots$ et s'il ne contient de facteurs de la forme $4m+3$ qu'élevés à des puissances paires, le nombre des solutions de l'équation $m^2 + n^2 = N$ est exprimé par

$$\rho_N = 4(a_1+1)(a_2+1)\dots;$$

mais si N contient un seul facteur premier de la forme $4m+3$, élevé à une puissance impaire, on aura $\rho_N = 0$.

Supposons d'abord que le nombre N ne contient que des facteurs premiers de la forme $4m+1$. Dans ce cas, la proposition résulte immédiatement du théorème trouvé; car le nombre des diviseurs du nombre N est $(a_1+1)(a_2+1) \dots$, et ils sont tous de la forme $4m+1$. Le nombre des diviseurs de cette forme ne varie pas quand N est multiplié par 2 ou une puissance de 2; mais si on le multiplie par un facteur premier de la forme $4m+3$, on introduit autant de diviseurs de cette forme que de la forme $4m+1$. Par conséquent on aura

$$a_N = b_N \text{ et } \rho_N = 0.$$

Si le nombre est de nouveau multiplié par le même facteur premier, on introduira de nouveau a_N facteurs de la forme $(4m+3)^2 = 4m'+1$, et ainsi de suite.

Voici comment la formule (11) peut être employée à la sommation de séries. Soit ρ_N le nombre de solutions de l'équation $m^2+n^2 = N$; on aura

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} f(m^2+n^2) = \sum_1^{\infty} \rho_N f(N),$$

où f est une fonction quelconque à laquelle correspond une série convergente. Nous supposons qu'on supprime dans la série double le terme unique correspondant à $m=n=0$, de sorte que N ne parcourt que les valeurs entières et positives.

On reconnaît facilement ici que $f(1)$, $f(2)$, ..., se trouvent le même nombre de fois dans les deux membres de l'équation, c'est-à-dire multipliés par le même facteur.

Si l'on se sert des notations employées ci-dessus, on aura de même

$$\sum_0^{\infty} \sum_1^{\infty} f((4m+1)n) = \sum_1^{\infty} a_N f(N),$$

$$\sum_0^{\infty} \sum_1^{\infty} f((4m+3)n) = \sum_1^{\infty} b_N f(N).$$

Comme $\rho_N = 4(a_N - b_N)$, on déduira de ces trois équations

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} f(m^2 + n^2)$$

$$= 4 \sum_0^{\infty} \sum_1^{\infty} [f((4m+1)n) - f((4m+3)n)]. \quad (13)$$

Par exemple

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} \frac{1}{(m^2 + n^2)^p}$$

$$= 4 \sum_0^{\infty} \sum_1^{\infty} \left[\frac{1}{(4m+1)^p n^p} - \frac{1}{(4m+3)^p n^p} \right],$$

où, comme nous l'avons dit, on n'a pas tenu compte dans le premier membre du terme où $m = n = 0$, et où l'exposant p est un nombre quelconque pour lequel la série est convergente. La dernière équation peut encore s'écrire

$$\frac{1}{4} \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} \frac{1}{(m^2 + n^2)^p}$$

$$= \left(1 + \frac{1}{2^p} + \frac{1}{3^p} + \dots\right) \left(1 - \frac{1}{3^p} + \frac{1}{5^p} - \dots\right).$$

La sommation double est ainsi ramenée à deux

sommations simples, qui peuvent être facilement exécutées avec une exactitude aussi grande qu'on veut.

Si p est un nombre entier, on peut encore ici, comme dans ce qui suit, effectuer avec exactitude l'une des sommations.

Nous pourrions de plus nous servir de la fraction $P(x)$ des équations (8) et (9) dans un autre cas. Posons $x = 0$ et nous aurons

$$\frac{2P(0)}{A_0} = \prod_1^{\infty} \left(\frac{1 - q^{2h}}{1 - q^{2h-1}} \right)^2 \frac{1 + q^{4h-2}}{1 + q^{4h}}.$$

Si l'on remarque que

$$1 - q^{2h} = \frac{1 - q^{4h}}{1 + q^{2h}} \quad \text{et} \quad 1 - q^{2h-1} = \frac{1 - q^{4h-2}}{1 + q^{2h-1}},$$

on reconnaîtra facilement que le produit infini ci-dessus coïncide avec le second membre de l'équation (7), où l'on a fait $c = 2$.

Des équations (7) et (9) on tirera par conséquent, pour $x = 0$ et $c = 2$,

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+2n^2} = 1 + 2 \left\{ \frac{q(1+q^2)}{1+q^4} + \frac{q^3(1+q^4)}{1+q^8} + \dots \right\} \quad (14);$$

le second membre pouvant s'écrire sous la forme

$$\begin{aligned} & 1 + 2 \sum_1^{\infty} \frac{q^n(1+q^{2n})}{1+q^{4n}} \\ &= 1 + 2 \sum_1^{\infty} \left[q^n - q^{3n} + q^{5n} - \dots + q^{3n} - q^{7n} + q^{11n} - \dots \right], \end{aligned}$$

on aura

$$\begin{aligned} & \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+2n^2} \\ &= 1 + 2 \sum_0^{\infty} \sum_1^{\infty} \left[q^{(8m+1)n} + q^{(8m+3)n} - q^{(8m+5)n} - q^{(8m+7)n} \right]. \quad (15) \end{aligned}$$

De ce résultat nous pouvons déduire le nombre de solutions de l'équation

$$m^2 + 2n^2 = N; \quad (16)$$

car, si l'on pose successivement

$$\begin{aligned} (8m+1)n &= N, & (8m+3)n &= N, \\ (8m+5)n &= N, & (8m+7)n &= N, \end{aligned}$$

et si l'on désigne les nombres de solutions de ces équations par a_N, b_N, c_N, d_N , où m et n ne peuvent prendre que des valeurs positives et entières, m à partir de 0 jusqu'à ∞ , n à partir de 1 jusqu'à ∞ , on aura

$$\rho_N = 2(a_N + b_N - c_N - d_N),$$

parce que N se trouve ce nombre de fois comme exposant dans les deux membres de l'équation (15). Mais, comme $a_N, b_N, \dots$, sont les nombres de diviseurs de N de la forme $(8m+1), (8m+3), \dots$, ils peuvent très facilement être déterminés.

De là on déduit la proposition suivante:

Si le nombre N contient les facteurs premiers $p_1, p_2, \dots$, de la forme $8m+1$ et $8m+3$ avec les exposants $\alpha_1, \alpha_2, \dots$, et si les facteurs premiers de la forme $8m+5$ et $8m+7$ n'y figurent qu'à des puissances paires, le nombre des solutions de l'équation $m^2 + 2n^2 = N$ sera exprimé par

$$\rho_N = 2(\alpha_1 + 1)(\alpha_2 + 1) \dots;$$

mais si un seul facteur premier de la forme $8m+5$ ou $8m+7$ entre dans N à une puissance impaire, on aura $\rho_N = 0$.

Pour parvenir à cette conclusion, nous pouvons nous servir d'un développement analogue à celui que

nous avons exposé relativement à l'équation $m^2 + n^2 = N$, si l'on remarque seulement que les puissances paires des nombres $8m+5$ et $8m+7$ sont nécessairement de la forme $8m+1$.

Si l'on emploie l'équation (15) à la sommation de séries, on obtiendra

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} f(m^2 + 2n^2) = 2 \sum_0^{\infty} \sum_1^{\infty} \left[f((8m+1)n) + f((8m+3)n) - f((8m+5)n) - f((8m+7)n) \right]. \quad (17)$$

Ici, comme dans l'équation (13), on n'a pas tenu compte du terme où $m = n = 0$ dans le premier membre.

Par exemple

$$\begin{aligned} & \frac{1}{2} \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} \frac{1}{[m^2 + 2n^2]^p} \\ &= \sum_1^{\infty} \frac{1}{n^p} \sum_0^{\infty} \left[\frac{1}{(8m+1)^p} + \frac{1}{(8m+3)^p} - \frac{1}{(8m+5)^p} - \frac{1}{(8m+7)^p} \right] \\ &= \left(1 + \frac{1}{2^p} + \frac{1}{3^p} + \dots \right) \left(1 + \frac{1}{3^p} - \frac{1}{5^p} - \frac{1}{7^p} + \frac{1}{9^p} + \frac{1}{11^p} - \dots \right). \end{aligned}$$

Il n'est pas difficile de trouver des fractions de forme autre que $P(x)$, qui peuvent être employées dans cette sorte de questions. Nous considérerons par exemple

$$\frac{1}{2-x} \prod_1^{\infty} \frac{1 - q^{2h-1}x + q^{4h-2}}{1 - q^{2h}x + q^{4h}} \cdot \frac{1 + q^{2h}x + q^{4h}}{1 + q^{2h-1}x + q^{4h-2}} = Q(x). \quad (18)$$

Cette fraction peut être décomposée en une série de fractions

$$Q(x) = \frac{A_0}{2-x} + \frac{A_1}{1+qx+q^2} + \frac{A_2}{1-q^2x+q^4} + \dots,$$

où nous trouverons, par les règles ordinaires de la décomposition des fractions rationnelles,

$$A_0 = \prod_1^{\infty} \left[\frac{(1 - q^{2h-1})(1 + q^{2h})^3}{(1 - q^{2h})(1 + q^{2h-1})} \right],$$

$$A_1 = 2qA_0, \quad A_2 = 2q^3A_0, \quad A_3 = 2q^5A_0, \quad \dots$$

Si l'on fait ici $x = 1$, on obtiendra

$$\frac{Q(1)}{A_0} = \prod_1^{\infty} \frac{1 - q^{2h-1} + q^{4h-2}}{1 - q^{2h} + q^{4h}} \cdot \frac{1 + q^{2h} + q^{4h}}{1 + q^{2h-1} + q^{4h-2}} \left[\frac{(1 - q^{2h})(1 + q^{2h-1})^3}{(1 - q^{2h-1})(1 + q^{2h})} \right]^3,$$

où nous pouvons poser

$$\begin{aligned} (1 - q^{2h-1} + q^{4h-2})(1 + q^{2h-1}) &= 1 + q^{6h-3}, \\ (1 + q^{2h} + q^{4h})(1 - q^{2h}) &= 1 - q^{6h}, \\ (1 + q^{2h-1} + q^{4h-2})(1 - q^{2h-1}) &= 1 - q^{6h-3}, \\ (1 - q^{2h} + q^{4h})(1 + q^{2h}) &= 1 + q^{6h}. \end{aligned}$$

Grâce à ces transformations, la série de facteurs ci-dessus devient identique au second membre de (7), où l'on a fait $c = 3$; et nous obtiendrons donc

$$\left. \begin{aligned} &\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2 + 3n^2} \\ &= 1 + \frac{2q}{1+q+q^3} + \frac{2q^3}{1-q^3+q^4} + \frac{2q^5}{1+q^3+q^5} + \dots \\ &= 1 + 2 \left[q \frac{1-q}{1-q^3} + q^3 \frac{1+q^3}{1+q^5} + q^5 \frac{1-q^3}{1-q^5} + \dots \right]. \end{aligned} \right\} \quad (19)$$

Si l'on se sert des signes de sommation, cette équation peut s'écrire sous la forme

$$\begin{aligned} &1 + 2 \sum_1^{\infty} \left[q^{2n-1} \frac{1 - q^{2n-1}}{1 - q^{6n-3}} + q^{2n} \frac{1 + q^{2n}}{1 + q^{6n}} \right] \\ &= 1 + 2 \sum_0^{\infty} \sum_1^{\infty} \left[q^{(8m+1)(2n-1)} (1 - q^{2n-1}) \right. \\ &\quad \left. + (q^{(8m+1)2n} - q^{(8m+4)2n}) (1 + q^{2n}) \right]. \end{aligned}$$

Dans cette identité entrent des termes de la forme $q^{(6m+1)2n}$ et $q^{(6m+5)2n}$, qui peuvent être éliminés au moyen des équations

$$\sum_0^{\infty} \sum_1^{\infty} q^{(6m+1)2n} = \sum_0^{\infty} \sum_1^{\infty} \left[q^{(3m+1)2n} - q^{(6m+4)2n} \right],$$

$$\sum_0^{\infty} \sum_1^{\infty} q^{(6m+5)2n} = \sum_0^{\infty} \sum_1^{\infty} \left[q^{(3m+2)2n} - q^{(6m+2)2n} \right],$$

dont la validité devient évidente, si l'on développe les deux membres suivant les valeurs croissantes de m . L'expression ci-dessus se réduira alors à

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+3n^2} = 1 + 2 \sum_0^{\infty} \sum_1^{\infty} \left[q^{(3m+1)n} - q^{(3m+2)n} \right. \\ \left. + 2q^{(3m+1)4n} - 2q^{(3m+2)4n} \right]. \quad (20)$$

En ce qui concerne les solutions de l'équation

$$m^2 + 3n^2 = N, \quad (21)$$

nous pouvons comme dans les cas précédents, en supposant que les équations

$$\begin{aligned} (3m+1)n &= N, & (3m+2)n &= N, \\ (3m+1)4n &= N, & (3m+2)4n &= N, \end{aligned}$$

aient respectivement a_N , b_N , c_N et d_N solutions en nombres positifs et entiers, conclure qu'on aura

$$\rho_N = 2(a_N - b_N) + 4(c_N - d_N).$$

De cette équation on peut déduire un théorème qui doit être considéré comme nouveau dans la théorie de nombres, parce qu'il ne peut pas être déduit immédiatement des théorèmes connus :

Si un nombre N contient les facteurs premiers $p_1, p_2, \dots$, de la forme $3m+1$ avec les exposants $\alpha_1, \alpha_2, \dots$ et si les facteurs premiers de la forme $3m+2$ n'y entrent qu'à des puissances paires, le nombre des solutions de l'équation $m^2 + 3n^2 = N$ est exprimé par

$$\rho_N = 2(\alpha_1 + 1)(\alpha_2 + 1) \dots,$$

si N est un nombre impair, et par

$$\rho_N = 6(\alpha_1 + 1)(\alpha_2 + 1) \dots$$

si N est pair. Si, au contraire, N contient un facteur premier de la forme $3m+2$ à une puissance impaire, on aura $\rho_N = 0$.

Si l'on applique l'équation (20) à la sommation de séries, on obtiendra

$$\left. \begin{aligned} & \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} f(m^2 + 3n^2) \\ = & 2 \sum_0^{\infty} \sum_1^{\infty} \left[f((3m+1)n) - f((3m+2)n) \right. \\ & \left. + 2f((3m+1)4n) - 2f((3m+2)4n) \right] \end{aligned} \right\} \quad (22)$$

où l'on n'a pas tenu compte du terme $m = 0$ dans le premier membre.

Par exemple

$$\begin{aligned} & \frac{1}{2} \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} \frac{1}{(m^2 + 3n^2)^p} \\ &= \sum_1^{\infty} \frac{1}{n^p} \sum_0^{\infty} \left[\frac{1}{(3m+1)^p} - \frac{1}{(3m+2)^p} \right] \left(1 + \frac{2}{4^p} \right) \\ &= \left(1 + \frac{1}{2^p} + \frac{1}{3^p} + \dots \right) \left(1 - \frac{1}{2^p} + \frac{1}{4^p} - \frac{1}{5^p} + \frac{1}{7^p} - \dots \right) \left(1 + \frac{2}{4^p} \right). \end{aligned}$$

Nous examinerons finalement la fraction

$$\frac{1}{2-x} \prod_{i=1}^{\infty} \frac{1+q^{2h}x+q^{4h}}{1+q^{2h-1}x+q^{4h-2}} \cdot \frac{1-q^{4h-2}x+q^{8h-4}}{1-q^{4h}x+q^{8h}} = R(x), \quad (23)$$

qui se décompose en une somme de fractions

$$\begin{aligned} \frac{A_0}{2-x} + \frac{A_1}{1+qx+q^2} + \frac{A_3}{1+q^3x+q^6} + \dots + \frac{A_4}{1-q^4x+q^8} \\ + \frac{A_8}{1-q^8x+q^{16}} + \dots, \end{aligned}$$

où les coefficients auront les valeurs

$$\begin{aligned} A_0 &= \prod_{i=1}^{\infty} \left[\frac{1+q^{2h}}{1+q^{2h-1}} \cdot \frac{1-q^{4h-2}}{1-q^{4h}} \right]^2, \\ A_1 &= qA_0, \quad A_3 = q^3A_0, \quad \dots, \quad A_4 = 2q^4A_0, \quad A_8 = 2q^8A_0, \quad \dots \end{aligned}$$

Si l'on fait ici $x = 0$, on trouvera

$$\frac{2R(0)}{A_0} = \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} q^{m^2+4n^2},$$

par où l'on peut obtenir un développement en série de cette somme double et, par des procédés analogues à ceux des cas précédents, une détermination du nombre des solutions de l'équation

$$m^2 + 4n^2 = N. \quad (24)$$

Ce résultat peut d'ailleurs être obtenu d'une autre manière plus simple, en comparant l'équation (24) ou bien $m^2 + (2n)^2 = N$ avec l'équation (12), $m^2 + n^2 = N$. On reconnaît facilement que, si N est ici un nombre impair, m sera dans le dernier cas impair, si n est pair et inversement, et que, par conséquent, le nombre des solutions de (24) est dans ce cas la moitié de celui de

l'équation (12). Si N est pair, mais non pas divisible par 4, m et n seront tous les deux impairs dans l'équation (12), et (24) n'aura pas de solutions. Finalement, si N est pair et divisible par 4, m et n seront tous les deux pairs dans l'équation (12), et (12) et (24) auront les mêmes solutions. Le résultat peut en tous cas être exprimé par la formule:

$$\sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} f(m^2 + 4n^2) = 2 \sum_0^{\infty} \sum_1^{\infty} \left[f((4m+1)n) - f((4m+3)n) - f((4m+1)2n) + f((4m+3)2n) + 2f((4m+1)4n) - 2f((4m+3)4n) \right] \quad (25)$$

par où l'on obtiendra, par exemple,

$$\begin{aligned} & \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} \frac{1}{(m^2 + 4n^2)^p} \\ &= \left(2 - \frac{1}{2^{p-1}} + \frac{1}{4^{p-1}} \right) \sum_1^{\infty} \frac{1}{n^p} \sum_0^{\infty} \left(\frac{1}{(4m+1)^p} - \frac{1}{(4m+3)^p} \right) \\ &= \left(\frac{1}{2} - \frac{1}{2^{p+1}} + \frac{1}{4^p} \right) \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} \frac{1}{(m^2 + n^2)^p}. \end{aligned}$$

Nous ferons encore en peu de mots mention de quelques autres cas dans lesquels on peut appliquer les formules développées.

Si dans l'équation (18) on fait $x = -2$ on obtiendra

$$\frac{4Q(-2)}{A_0} = \prod_1^{\infty} \left[\frac{1 - q^{2h}}{1 + q^{2h}} \cdot \frac{1 + q^{2h-1}}{1 - q^{2h-1}} \right]^4.$$

Ce produit infini est d'après (6) égal à

$$\left[\sum_{-\infty}^{+\infty} q^{n^2} \right]^4 = \sum_{-\infty}^{+\infty} q^{7(4)} q^{7^2 + m^2 + n^2 + p^2},$$

où le signe de sommation du second membre désigne une sommation quadruple étendue à toutes les valeurs entières des quatres nombres l, m, n, p à partir de $-\infty$ jusqu'à $+\infty$, y compris zéro. Pour cette somme quadruple on obtiendra le développement en série connu de la théorie des fonctions elliptiques et par conséquent la détermination du nombre de solutions de l'équation

$$l^2 + m^2 + n^2 + p^2 = N$$

et la formule de sommation pour $\Sigma^{(4)} f(l^2 + m^2 + n^2 + p^2)$.

De plus, si dans l'équation (23) on pose $x = -2$, on parviendra au cas de l'équation

$$l^2 + m^2 + 2n^2 + 2p^2 = N.$$

J'ajouterai les résultats suivants, dont on peut se

* NOTE. servir quelquefois en physique mathématique*

$$\begin{aligned} & \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} f(m^2 + n^2 + mn) \\ &= 6 \sum_0^{\infty} \sum_1^{\infty} \left[f((3m+1)n) - f((3m+2)n) \right], * \end{aligned}$$

où le terme $n = m = 0$ doit être négligé, et

$$\begin{aligned} & \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} f(3(m^2 + n^2 + mn + m + n) + 1) \\ & 3 \sum_0^{\infty} \sum_1^{\infty} \left[f((3m+1)(6n-5)) - f((3m+2)(6n-1)) \right], \end{aligned}$$

* La démonstration de ce théorème est développée dans un mémoire, que j'ai l'intention de publier dans le journal de Borchardt.

où l'on doit, dans le premier membre, tenir compte de tous les termes.

Si, dans l'équation $m^2 + cn^2 = N$, c est négatif, le cas diffère essentiellement des cas traités ci-dessus. Dans ce cas, le nombre de solutions est en général infini, sauf dans le seul cas où $-c$ est un carré parfait.

Nous considérerons ici le seul cas de $c = -1$ ou

$$m^2 - n^2 = N. \quad (26)$$

La méthode analytique consiste ici encore à transformer la série

$$\sum_{-\infty}^{+\infty} \sum q^{m^2 - n^2};$$

mais comme l'exposant est égal à N , nombre entier et positif différent de 0, on doit ici effectuer la sommation de manière que $m^2 - n^2$ reste toujours positif. Par suite n ne peut prendre que les valeurs entières, positives et négatives, qui sont numériquement plus petites que m , tandis que m peut recevoir toutes les valeurs entières de $-\infty$ jusqu'à $+\infty$, 0 seul étant excepté.

Par conséquent la série sera

$$2[q + q^4 + 2q^5 + q^9 + 2q^8 + 2q^5 + q^{16} + 2q^{15} + 2q^{12} + 2q^7 + \dots],$$

où la série entre crochets peut être décomposée de la manière suivante

$$q + 2q^3 + 2q^5 \dots = q \frac{1 + q^2}{1 - q^2},$$

$$q^4 + 2q^8 + 2q^{12} \dots = q^4 \frac{1 + q^4}{1 - q^4},$$

$$q^9 + 2q^{15} + 2q^{21} \dots = q^9 \frac{1 + q^6}{1 - q^6},$$

et ainsi de suite; nous obtiendrons de cette manière

$$\frac{1}{2} \sum_{-\infty}^{+\infty} \sum q^{m^2-n^2} = \sum_1^{\infty} q^{n^2} \frac{1+q^{2n}}{1-q^{2n}}.$$

Le second membre de cette équation peut être remplacé par les deux sommes

$$\sum_1^{\infty} q^{4n^2} \frac{1+q^{4n}}{1-q^{4n}} \quad \text{et} \quad \sum_1^{\infty} q^{(2n-1)^2} \frac{1+q^{4n-2}}{1-q^{4n-2}}$$

que nous traiterons séparément.

La première peut être rapportée à la forme générale

$$\sum_1^{\infty} q^{n^2} x^{2n-1} \frac{1+q^n x}{1-q^n x} = \varphi(x), \quad (28)$$

expression qui sera identique à la somme ci-dessus, si l'on fait $x = 1$ et si l'on remplace q par q^4 .

De la dernière équation ou

$$\varphi(x) = qx \frac{1+qx}{1-qx} + q^4 x^3 \frac{1+q^2 x}{1-q^2 x} + q^9 x^5 \frac{1+q^3 x}{1-q^3 x} + \dots$$

on tirera

$$q^2 x^2 \varphi(qx) = q^4 x^3 \frac{1+q^2 x}{1-q^2 x} + q^9 x^5 \frac{1+q^3 x}{1-q^3 x} + \dots,$$

et, par conséquent, l'équation fondamentale

$$\varphi(x) - q^2 x^2 \varphi(qx) = qx \frac{1+qx}{1-qx}. \quad (29)$$

$\varphi(x)$ est complètement déterminée par cette équation, car le développement en série par lequel $\varphi(x)$ est primitivement définie peut facilement être déduit de cette équation. Par conséquent toute fonction qui satisfait à

l'équation fondamentale sera identique à $\varphi(x)$. Mais il est évident que la série

$$\frac{qx}{1-qx} + \frac{q^2x^2}{1-q^2x} + \frac{q^3x^3}{1-q^3x} + \dots = \varphi(x), \quad (qx < 1)$$

satisfait à cette équation fondamentale, car on aura

$$q^2x^2\varphi(qx) = \frac{q^4x^3}{1-q^2x} + \frac{q^6x^4}{1-q^3x} + \frac{q^8x^5}{1-q^4x} + \dots,$$

et par conséquent

$$\begin{aligned} \varphi(x) - q^2x^2\varphi(qx) &= \frac{qx}{1-qx} + q^2x^2 + q^3x^3 + q^4x^4 + \dots \\ &= qx \frac{1+qx}{1-qx}. \end{aligned}$$

Nous parviendrons de cette manière à l'équation remarquable

$$\sum_1^{\infty} q^{n^2} x^{2n-1} \frac{1+q^n x}{1-q^n x} = \sum_1^{\infty} \frac{q^n x^n}{1-q^n x}. \quad (30)$$

Si l'on fait ici $x = 1$ et qu'on remplace q par q^4 , on obtiendra

$$\sum_1^{\infty} q^{4n^2} \frac{1+q^{4n}}{1-q^{4n}} = \sum_1^{\infty} \frac{q^{4n}}{1-q^{4n}}. \quad (31)$$

La seconde somme qui doit être transformée peut être rapportée à l'expression

$$\sum_1^{\infty} q^{(2n-1)^2} x^{2n-1} \frac{1+q^{4n-2}x}{1-q^{4n-2}x} = \psi(x), \quad (32)$$

fonction qui, comme on le reconnaît facilement, satisfait à l'équation fondamentale

$$\psi(x) - q^4x^2\psi(q^4x) = qx \frac{1+q^2x}{1-q^2x}. \quad (33)$$

Par là $\phi(x)$ est complètement déterminée. La même équation est vérifiée par la série

$$\frac{qx}{1-q^2x} + \frac{q^3x^3}{1-q^6x} + \frac{q^5x^5}{1-q^{10}x} + \dots = \phi(x),$$

car cette série donne pour, $q^2x < 1$,

$$\begin{aligned}\phi(x) - q^4x^3\phi(q^4x) &= \frac{qx}{1-q^2x} + q^3x^3 + q^5x^5 + q^7x^7 + \dots \\ &= qx \frac{1+q^2x}{1-q^2x}.\end{aligned}$$

On aura donc

$$\sum_1^{\infty} q^{(2n-1)^2} x^{2n-1} \frac{1+q^{4n-2}x}{1-q^{4n-2}x} = \sum_1^{\infty} \frac{q^{2n-1}x^n}{1-q^{4n-2}x}; \quad (34)$$

d'où, pour $x = 1$,

$$\sum_1^{\infty} q^{(2n-1)^2} \frac{1+q^{4n-2}}{1-q^{4n-2}} = \sum_1^{\infty} \frac{q^{2n-1}}{1-q^{4n-2}}. \quad (35)$$

En additionnant les deux séries (31) et (35), on obtiendra la transformation en question des séries qui entrent dans l'équation (27), savoir

$$\frac{1}{2} \sum_{-\infty}^{+\infty} \sum q^{m^2-n^2} = \sum_1^{\infty} \left[\frac{q^{4n}}{1-q^{4n}} + \frac{q^{2n-1}}{1-q^{4n-2}} \right],$$

ou bien, en mettant aussi la dernière série sous la forme d'une somme double,

$$\frac{1}{2} \sum_{-\infty}^{+\infty} \sum q^{m^2-n^2} = \sum_1^{\infty} \sum_1^{\infty} \left[q^{4mn} + q^{(2m-1)(2n-1)} \right]. \quad (36)$$

On peut conclure de là que la moitié du nombre des solutions de l'équation $m^2 - n^2 = N$ est égale au

nombre des solutions (m et n positifs) de l'équation

$$4mn = N \text{ ou de l'équation } (2m-1)(2n-1) = N;$$

ces deux équations ne peuvent pas avoir lieu en même temps. Ce résultat peut encore être exprimé ainsi :

Le nombre des solutions de l'équation $m^2 - n^2 = N$ est égal au double du nombre des diviseurs de N ou de $\frac{1}{4}N$, selon que N est impair ou divisible par 4. Si, au contraire, N n'est divisible que par 2, l'équation n'a point de solutions.

En conséquence de (36), on aura de plus

$$\begin{aligned} & \frac{1}{2} \sum_{-\infty}^{+\infty} f(m^2 - n^2) \\ &= \sum_1^{\infty} \sum_1^{\infty} [f(4mn) + f(2m-1)(2n-1)]. \quad (37) \end{aligned}$$

Par exemple,

$$\begin{aligned} & \frac{1}{2} \sum_{-\infty}^{+\infty} \sum \frac{1}{(m^2 - n^2)^p} \\ &= \frac{1}{4^p} \sum_1^{\infty} \frac{1}{m^p} \sum_1^{\infty} \frac{1}{n^p} + \sum_1^{\infty} \frac{1}{(2m-1)^p} \sum_1^{\infty} \frac{1}{(2n-1)^p} \\ &= \left(\frac{1}{2^p} + \frac{1}{4^p} + \dots \right)^2 + \left(1 + \frac{1}{3^p} + \frac{1}{5^p} + \dots \right)^2. \end{aligned}$$

Si p est un nombre pair, on peut trouver la valeur exacte de ces deux sommes.

NOTE.

Lorenz n'a jamais publié aucun mémoire sur ce sujet dans le journal de Borchardt. Mais les deux théorèmes qu'il cite peuvent facilement être déduits du dernier théorème démontré pag. 421 sur le nombre des solutions de l'équation $N = m^2 + 3n^2$.

L'équation

$$N = m^2 + n^2 + mn \quad (a)$$

peut s'écrire

$$4N = (2m + n)^2 + 3n^2, \quad (b)$$

par où l'on reconnaît immédiatement que toute solution de l'équation en question donne une solution de l'équation

$$4N = m_1^2 + 3n_1^2, \quad (c)$$

et comme, dans cette équation, m_1 et n_1 sont en même temps pairs ou impairs, toute solution de l'équation (c) fournira une solution de l'équation (a). Le nombre des solutions de (a) est donc égal à celui de (c), ce qui démontre la première proposition de Lorenz.

La seconde peut être démontrée d'une manière analogue. Si l'on a

$$N = 3(m^2 + n^2 + mn + m + n) + 1, \quad (d)$$

on aura

$$4N = 3(2m + n + 1)^2 + (3n + 1)^2. \quad (e)$$

Mais N et par conséquent $4N$ doivent être de la forme $3p+1$, et si l'on pose

$$4N = 3m_1^2 + n_1^2, \quad (f)$$

n_1 sera de la forme $3p \pm 1$; par conséquent, soit n_1 soit $-n_1$ qui tous deux satisfont à l'équation (f), sera de la forme $3p+1$, tandis que m_1 et n_1 seront en même temps pairs ou impairs. Par conséquent, à toute solution de (f) pour laquelle n_1 est de la forme $3p+1$ correspondra une solution de l'équation (e). L'équation (e) a donc un nombre de solutions qui est la moitié de celui de l'équation (f).

Par conséquent, on aura

$$\begin{aligned} & \sum_{-\infty}^{+\infty} \sum_{-\infty}^{+\infty} f(3(m^2+n^2+mn+m+n)+1) \\ &= 3 \sum_0^{\infty} \sum_0^{\infty} [f((3m+1)(3n+1)) - f((3m+2)(3n+2))]. \end{aligned}$$

De plus, on peut négliger toutes les valeurs de $3n+1$ et $3n+2$ qui sont paires; car à toute solution de l'équation

$$N = (3m+1)(3n+1),$$

où $(3n+1)$ est un nombre pair, correspond une solution de l'équation

$$N = (3m_1+2)(3n_1+2),$$

à savoir

$$N = (6m+2) \left(\frac{3n+1}{2} \right),$$

où $(6m+2)$ est pair.

Mais un nombre impair de la forme $3p+1$ peut toujours être écrit sous la forme $6n-5$, et un nombre impair de la forme $3p+2$ sous la forme $6n-1$.

SUR LA COMPENSATION
DES ERREURS D'OBSERVATION.

SUR LA COMPENSATION DES ERREURS D'OBSERVATION.

TIDSSKRIFT FOR MATEMATIK, 1872, P. 1—20.*

* NOTE 1.

(Lu dans la séance de l'Académie des sciences, à Copenhague, le 12 janvier 1872.)

I.

Introduction.

Si l'on a plusieurs observations qui dépendent en quelque manière les unes des autres, on peut souvent déterminer les valeurs particulières observées avec une précision plus grande que celle qui ressort des observations immédiates. La condition pour qu'on puisse de cette manière diminuer ou compenser les erreurs commises dans les observations consiste en général en ceci, que les vraies valeurs des quantités observées soient des fonctions d'une ou plusieurs variables, dont les valeurs vraies ou probables sont connues pour chaque observation.

Si ces fonctions diffèrent pour toutes les observations, on doit, au moins, quand une compensation est possible, connaître leurs formes, et le nombre des constantes inconnues dont elles dépendent doit être plus petit que celui des observations. Au contraire, si les fonctions elles-mêmes sont complètement inconnues, ce qui est souvent le cas en réalité, il faut qu'on sache que les

valeurs vraies de plusieurs quantités observées peuvent être exprimées par la même fonction.

Les valeurs observées, qui, du reste, au lieu d'être immédiatement observées, peuvent être le résultat d'une ou plusieurs observations, combinées avec des quantités indépendantes des observations, étant désignées par

$$O_1, O_2, O_3, \dots, O_n,$$

les valeurs vraies correspondantes seront représentées par

$$O'_1, O'_2, O'_3, \dots O'_n.$$

Si ces dernières quantités sont des fonctions d'une ou plusieurs variables, on pourra, au moins approximativement, exprimer ces quantités comme fonctions linéaires de constantes inconnues ou éléments, $\mathcal{A}', \mathcal{B}', \mathcal{C}' \dots$, dont le nombre est e , de manière qu'on peut approximativement ou exactement écrire

$$\left. \begin{aligned} O'_1 &= a_1 \mathcal{A}' + b_1 \mathcal{B}' + c_1 \mathcal{C}' \dots \\ O'_2 &= a_2 \mathcal{A}' + b_2 \mathcal{B}' + c_2 \mathcal{C}' \dots \\ O'_3 &= a_3 \mathcal{A}' + b_3 \mathcal{B}' + c_3 \mathcal{C}' \dots \\ &\dots \dots \dots \end{aligned} \right\} \quad (1)$$

où maintenant les coefficients $a, b, c, \dots$ affectés d'indices sont des fonctions connues ou inconnues des variables connues pour chaque observation.

S'il n'y a, par exemple, qu'une quantité observée, on aura $O'_1 = O'_2 = O'_3 \dots$, et l'on pourra poser $a_1 = a_2 = a_3 \dots = 1$, tandis que les autres termes s'évanouiront, de manière que $\mathcal{A}'$ devient la seule quantité inconnue dont la valeur probable doit être déterminée par les observations.

S'il n'y a, pour citer un autre exemple très étendu, qu'une variable indépendante t , qui pour les différentes

observations $O_1, O_2, O_3 \dots$ a les valeurs $t_1, t_2, t_3 \dots$, $a_1, b_1, c_1 \dots$ seront des fonctions de $t_1, a_2, b_2, c_2 \dots$ des fonctions de $t_2 \dots$. Si, de plus, ces fonctions ne sont pas connues, mais si l'on sait que $O'_1, O'_2, O'_3 \dots$ sont exprimées respectivement par la même fonction de $t_1, t_2, t_3 \dots$, on pourra essayer d'obtenir des déterminations de $O'_1, O'_2, O'_3 \dots$ soit par un développement en série suivant les puissances de t , par où on aura $a_1 = 1$, $b_1 = t_1$, $c_1 = t_1^2$, $\dots$ et ainsi de suite, soit par un développement suivant les puissances d'une fonction de t appropriée aux circonstances données. Parfois un développement suivant des fonctions périodiques peut être avantageux.

Il est un cas où les coefficients $a, b, c \dots$ peuvent être choisis arbitrairement; c'est celui où l'on suppose que le nombre des éléments des équations (1) est égal au nombre n des observations. Mais, dans ce cas, on ne pourrait évidemment obtenir aucune compensation au moyen de ces équations, et elles ne fourniraient aucune relation entre les observations.

Le dernier exemple indique déjà une proposition qui résultera avec plus de précision des recherches suivantes, à savoir que *les valeurs $a, b, c \dots$ qui satisfont exactement aux équations (1), et que nous dénommerons les vraies valeurs, ne sont pas nécessairement identiques à celles qui donnent la compensation la plus avantageuse.* Le seul cas d'exception est celui où il n'y a qu'un élément.

On a toujours jusqu'ici, en exécutant une compensation au moyen de la méthode des moindres carrés, considéré les valeurs vraies des coefficients $a, b, c \dots$ comme des quantités données, dont les valeurs devraient

nécessairement passer sans changement des expressions vraies des observations aux valeurs compensées. On a ainsi négligé l'existence possible d'autres valeurs de ces coefficients qui pourraient amener une compensation plus avantageuse, et l'on a écarté le cas où il serait, par exemple, préférable de supprimer l'une d'elles que de la conserver. Les conditions sous lesquelles une telle omission doit être considérée comme plus avantageuse seront précisément indiquées dans ce qui suivra.

Mais si, même dans le cas où l'on connaît les valeurs vraies des coefficients a , b , c , ..., la question n'est pas tranchée, de savoir s'il n'y a pas d'autres valeurs plus favorables pour la compensation, il y a d'autant plus lieu de poser cette question dans le cas où l'on ne connaît pas les valeurs de ces coefficients, mais où ils ne peuvent être exprimés qu'approximativement, par exemple, par un développement en série. Le problème à résoudre sera donc de trouver combien de termes de la série et lesquels on doit conserver, puis de décider, si l'on a essayé de se servir de différents développements, lequel d'entre eux donne la compensation la plus avantageuse.

Nous ne pouvons donc dans aucun cas, sauf dans le cas d'un élément unique, considérer les coefficients a , b , c , ..., qui seront employés par la compensation, comme des données. Mais, d'autre part, il ne peut être question de déterminer au moyen *des quantités observées* les fonctions des variables indépendantes qui donneraient absolument la compensation la plus favorable. Le problème doit être limité de manière à indiquer les conditions générales de la compensation la plus favorable sous une forme telle qu'il soit toujours possible de

choisir, parmi plusieurs compensations qui pourraient être employées, celle qui doit être considérée comme la plus favorable.

2.

Transformation des coefficients.

Dans ce qui suit on se servira souvent d'expressions de la forme

$$a_r \mathcal{A} + b_r \mathcal{B} + c_r \mathcal{C} + d_r \mathcal{D} + \dots,$$

où l'indice r représente l'un quelconque des nombres 1, 2, 3 ... n . Il sera convenable de transformer d'abord ces expressions par l'introduction des quantités nouvelles $\alpha_1, \alpha_2, \beta_1, \beta_2, \gamma_1, \dots$, qui sont déterminées par les équations suivantes

$$\left. \begin{aligned} [pa\alpha] \alpha_1 + [pa\beta] \beta_1 &= 0, \\ [pa\alpha] \alpha_2 + [pa\beta] \beta_1 + [pac] &= 0, \\ [pba] \alpha_2 + [pbb] \beta_1 + [pbc] &= 0, \\ [pa\alpha] \alpha_3 + [pa\beta] \beta_2 + [pac] \gamma_1 + [pab] &= 0, \\ [pba] \alpha_3 + [pbb] \beta_2 + [pbc] \gamma_1 + [pbd] &= 0, \\ [pca] \alpha_3 + [pcb] \beta_2 + [pcc] \gamma_1 + [pcb] &= 0, \\ \dots \dots \dots \end{aligned} \right\} \quad (2)$$

où l'on s'est servi de la notation abrégée ordinaire de la sommation, à savoir $[pab]$ au lieu de $p_1 \alpha_1 b_1 + p_2 \alpha_2 b_2 + \dots + p_n \alpha_n b_n$ etc. Les quantités $p_1, p_2, \dots, p_n$ désigneront dans ce qui suit les poids correspondants aux observations $O_1, O_2, \dots, O_n$.

Puis, si l'on pose

$$\left. \begin{aligned} a_r &= a_r, \\ b_r &= a_r \alpha_1 + b_r, \\ c_r &= a_r \alpha_2 + b_r \beta_1 + c_r, \\ d_r &= a_r \alpha_3 + b_r \beta_2 + c_r \gamma_1 + d_r, \\ \dots \dots \dots \end{aligned} \right\} \quad (3)$$

on aura

$$\left. \begin{aligned} [paa] &= [pa\alpha], \\ [pbb] &= [pa\beta] \alpha_1 + [p\beta\beta], \\ [pcc] &= [pa\gamma] \alpha_2 + [p\beta\gamma] \beta_1 + [p\gamma\gamma], \\ &\dots \dots \dots \end{aligned} \right\} \quad (4)$$

tandis que les équations (2) donneront

$$\left. \begin{aligned} [pab] &= 0, & [pac] &= 0, & [pad] &= 0, & \dots \\ & & [pbc] &= 0, & [pbd] &= 0, & \dots \\ & & [pcd] &= 0, & & & \dots \end{aligned} \right\} \quad (5)$$

Si nous introduisons de plus, dans l'expression considérée

$$a_r \mathcal{A} + b_r \mathcal{B} + c_r \mathcal{C} + d_r \mathcal{D} + \dots,$$

les constantes nouvelles $A, B, C, \dots$, au moyen des équations

$$\left. \begin{aligned} \mathcal{A} &= A + Ba_1 + Ca_2 + Da_3 + \dots, \\ \mathcal{B} &= \quad B + C\beta_1 + D\beta_2 + \dots, \\ \mathcal{C} &= \quad \quad C + D\gamma_1 + \dots, \end{aligned} \right\} \quad (6)$$

l'expression sera transformée en une autre de la même forme et ayant le même nombre de termes, savoir

$$a_r A + b_r B + c_r C + d_r D + \dots \quad (7)$$

Dans cette expression les nouveaux coefficients $a, b, c, \dots$ satisferont donc aux équations (5).

3.

Compensation par des fonctions linéaires des observations.

Les quantités compensées qui correspondent aux quantités observées

$$O_1, \quad O_2, \quad O_3, \quad \dots, \quad O_n$$

seront désignées par

$$O_1 + v_1, \quad O_2 + v_2, \quad O_3 + v_3, \quad \dots, \quad O_n + v_n.$$

On peut toujours supposer que le calcul de ces valeurs soit exécuté au moyen d'expressions de la forme (7), de manière qu'on peut poser

$$O_r + v_r = a_r A + b_r B + c_r C + \dots \quad (8)$$

où $r = 1, 2, \dots, n$. Les coefficients $a, b, c, \dots$ sont formés avec $\alpha, \beta, \gamma, \dots$ de la manière indiquée dans le paragraphe précédent et ils satisferont donc aux équations (5).

Comme les coefficients originaires $\alpha, \beta, \gamma, \dots$ ne doivent pas nécessairement satisfaire exactement aux équations (1), nous pourrons, en introduisant les nouveaux coefficients et les nouvelles constantes $A', B', C', \dots$, écrire

$$O'_r + v'_r = a_r A' + b_r B' + c_r C' + \dots \quad (9)$$

Tandis qu'on doit de cette manière, si l'on se sert des coefficients choisis $\alpha, \beta, \gamma, \dots$ introduire des corrections (v'_r) dans les valeurs vraies (O'_r) correspondantes aux observations, nous pourrons imaginer une troisième suite de quantités

$$O''_1, \quad O''_2, \quad O''_3, \quad \dots, \quad O''_n,$$

qui à présent se comporteront comme les valeurs vraies, c'est-à-dire comme des quantités pour lesquelles les corrections calculées s'évanouiront. Nous aurons donc pour cette suite

$$O''_r = a_r A'' + b_r B'' + c_r C'' + \dots \quad (10)$$

Les constantes $A, B, C, \dots$, des équations (8) doivent en général être considérées comme des fonctions des valeurs observées et des coefficients choisis $a, b, c, \dots$

Il sera pourtant nécessaire de restreindre ici la généralité du problème par la condition que $A, B, C, \dots$ soient des fonctions linéaires des valeurs observées.

Nous poserons donc

$$\left. \begin{aligned} A &= x_1 O_1 + x_2 O_2 + \dots + x_n O_n = [xO], \\ B &= [yO], \quad C = [zO], \quad \dots, \end{aligned} \right\} \quad (11)$$

où $x, y, z, \dots$, affectés d'indices dépendent des coefficients $a, b, c, \dots$ et par conséquent aussi de $a, b, c, \dots$ mais sont indépendants des valeurs observées.

L'équation (8) est ainsi transformée en

$$O_r + v_r = a_r [xO] + b_r [yO] + c_r [zO] \dots \quad (12)$$

Si les quantités O_r'' sont calculées de la même manière, on aura

$$A'' = [xO''], \quad B'' = [yO''], \quad C'' = [zO''], \quad \dots; \quad (13)$$

mais les corrections calculées doivent ici s'évanouir, et par conséquent les équations (10) seront exactement vérifiées par ces valeurs des constantes. C'est pourquoi nous pouvons éliminer les quantités O_r'' . Car, si l'on multiplie l'équation (10) par x_r , et si l'on additionne les différentes expressions obtenues en posant $r = 1, 2, \dots, n$, on aura

$$[xO''] = [ax] A'' + [bx] B'' + [cx] C'' + \dots,$$

et par conséquent

$$A'' = [ax] A'' + [bx] B'' + [cx] C'' + \dots$$

et par analogie

$$\begin{aligned}
B'' &= [ay] A'' + [by] B'' + [cy] C'' + \dots \\
C'' &= [az] A'' + [bz] B'' + [cz] C'' + \dots \\
&\dots \dots \dots
\end{aligned}$$

Cependant, comme les coefficients $A'', B'', C'', \dots$ sont des quantités qui ne peuvent pas être déterminées, les équations formées de cette manière doivent être identiques, et les coefficients x, y, z doivent donc satisfaire aux équations

$$\left. \begin{aligned}
[ax] &= 1, & [ay] &= 0, & [az] &= 0, & \dots \\
[bx] &= 0, & [by] &= 1, & [bz] &= 0, & \dots \\
[cx] &= 0, & [cy] &= 0, & [cz] &= 1, & \dots \\
&\dots \dots \dots
\end{aligned} \right\} \quad (14)$$

Le nombre de ces équations est e^2 , si e est le nombre des éléments, tandis que le nombre des coefficients $x, y, z, \dots$ est en . C'est seulement dans le cas de $e = n$ que les coefficients $x, y, z, \dots$ seraient complètement déterminés par les équations (14), et, si l'on employait les valeurs ainsi déterminées, on trouverait toutes les corrections $v_r = 0$.

Si les coefficients $x, y, z, \dots$ sont remplacés par les coefficients nouveaux $\xi, \eta, \zeta, \dots$ définis par les équations

$$\left. \begin{aligned}
x_r &= \frac{p_r a_r}{[p a a]} + \xi_r, & y_r &= \frac{p_r b_r}{[p b b]} + \eta_r, \\
z_r &= \frac{p_r c_r}{[p c c]} + \zeta_r + \dots,
\end{aligned} \right\} \quad (15)$$

les équations de condition de ces nouveaux coefficients seront

$$\left. \begin{aligned}
[a\xi] &= 0, & [a\eta] &= 0, & [a\zeta] &= 0, & \dots \\
[b\xi] &= 0, & [b\eta] &= 0, & [b\zeta] &= 0, & \dots \\
&\dots \dots \dots
\end{aligned} \right\} \quad (16)$$

Les conditions (14) ou (16) doivent donc être remplies pour toute compensation dans laquelle les valeurs compensées des observations doivent être des fonctions linéaires des valeurs observées. Reste à déterminer les conditions de la compensation la plus favorable; mais cette question ne peut être tranchée que par un calcul de probabilités, et les règles qui doivent être employées seront traitées en détail dans le paragraphe suivant.

4.

Calculs de probabilités.

Après qu'on a exécuté une observation, l'erreur qu'on a commise sera une quantité complètement déterminée, savoir la différence des valeurs vraie et observée; mais *avant* l'observation et, par conséquent, avant que cette quantité ait reçu une valeur déterminée, il y a *la possibilité* d'une multitude infinie de valeurs pour l'observation et pour son erreur. Nous partirons de ce point de vue aprioristique, en supposant que le calcul est exécuté avant toute observation et qu'on doit, par suite, tenir compte de toutes les valeurs possibles de l'erreur.

En ce qui concerne les erreurs d'observation elles-mêmes, nous supposerons que la probabilité des erreurs possibles de chaque observation prise séparément est déterminée par la loi exponentielle des erreurs. La validité de cette supposition est, comme on sait, fondée sur ce que l'on peut imaginer que les erreurs sont produites par une multitude d'erreurs partielles indépendantes l'une de l'autre qui, prises chacune séparément, influent sur l'erreur totale. Cela implique une supposition double en ce qui concerne les observations elles-mêmes, à savoir:

1) Que l'erreur commise dans une observation particulière n'influe sur aucune autre observation;

2) Que chaque erreur d'observation peut indifféremment prendre les mêmes valeurs positives et négatives.

Outre les erreurs d'observation, une autre espèce d'erreurs intervient encore dans les valeurs compensées; c'est celle qui provient de la différence entre les coefficients choisis $a, b, c, \dots$, et leurs valeurs vraies. Tandis que les erreurs d'observation, que nous désignerons par u , sont déterminées par

$$u_r = O'_r - O_r, \quad (17)$$

l'autre espèce d'erreurs, que nous nommerons *erreurs de formule* et que nous désignerons par u' , peut être définie par l'équation

$$u'_r = O''_r - O'_r. \quad (18)$$

Pour ces erreurs on aura de même *avant* le choix des coefficients $a, b, c, \dots$, une infinité de valeurs possibles; mais nous n'avons aucun moyen de déterminer les lois des erreurs valables ici, et surtout il n'y a aucune raison pour supposer ici la validité de la loi exponentielle, comme nous l'avons fait pour les erreurs d'observation.

Nous pourrions pourtant toujours supposer remplies les deux conditions suivantes: 1) toutes les erreurs de formule sont indépendantes l'une de l'autre; 2) elles peuvent toutes prendre indifféremment les mêmes valeurs positives et négatives. Or les n erreurs de formule ne dépendent que du choix des n coefficients $a, b, c, \dots$, et l'on peut toujours supposer que les choix arbitraires, en nombre infini, de ces coefficients varient de manière

à faire qu'une erreur de formule parcoure toute la série de ses valeurs, tandis que les autres ne varient pas et qu'en même temps chacune des erreurs prend aussi souvent les mêmes valeurs positives que négatives. Il va sans dire que les erreurs de formule sont en tout cas indépendantes des erreurs dues aux observations elles-mêmes.

Nous définissons comme à l'ordinaire „la valeur moyenne“ d'une fonction des erreurs comme la somme des valeurs que la fonction prend en tous cas possibles, divisée par le nombre de ces cas. Si la même valeur se présente plusieurs fois, il faut la compter autant le fois qu'elle intervient. Comme notation de la moyenne d'une fonction on se servira dans ce qui suit d'un trait horizontal sur le symbole de fonction: ainsi $\overline{u}$ désigne la moyenne de u . Il suit de nos suppositions sur les erreurs que la moyenne d'un produit de deux facteurs dont l'un est une fonction d'erreurs qui n'entrent pas dans l'autre est égale au produit des moyennes des deux facteurs, et que la moyenne d'une puissance impaire d'une erreur est nulle. Si nous remplaçons une fonction U des erreurs inconnues, par V , l'erreur commise sera $U - V$. „L'erreur moyenne“ M de la même supposition sera définie par

$$M^2 = (\overline{U - V})^2.$$

Plus petite est cette erreur moyenne, plus „avantageuse“ est jugée la détermination. Comme M^2 a sa plus petite valeur pour $V = \overline{U}$, savoir la valeur

$$M^2 = \overline{U^2} - (\overline{U})^2, \quad (19)$$

la plus avantageuse détermination de u_r sera $u_r = \overline{u_r}$

$= 0$, et le carré de son erreur moyenne sera $\overline{u_r^2}$. Le „poids“ p de cette détermination ou de la valeur observée O_r est inversement proportionnel au carré de l'erreur moyenne de la détermination, d'où il suit que

$$p_1 \overline{u_1^2} = p_2 \overline{u_2^2} \dots = p_n \overline{u_n^2} = m^2, \quad (20)$$

où m désigne l'erreur moyenne de l'unité arbitraire de poids. Ces relations ne sont valables que pour les erreurs d'observation, mais non pour les erreurs de formule.

Comme nous supposons que les erreurs d'observation sont assujetties à la loi exponentielle des erreurs, nous dirons, conformément au principe d'où la „méthode des moindres carrés“ a tiré son nom, que les valeurs les plus probables des valeurs observées seront celles pour lesquelles la somme

$$[p(O' - O - v)^2] = [p(u - v)^2]^* \quad \text{* NOTE 2.}$$

devient un minimum. Dans la méthode ordinairement suivie pour le calcul des valeurs les plus probables, on n'a pas considéré les coefficients $a, b, c, \dots$, comme variables, et le problème peut alors, comme on sait, être réduit à chercher le minimum de la somme $[pv^2]$. Mais on reconnaît facilement que cela ne saurait être permis ici, où les coefficients $a, b, c, \dots$, sont supposés variables, par cette première raison que, dans le cas où l'on choisirait autant d'éléments que d'observations, toutes les corrections v_r s'évanouiraient et que par conséquent le minimum absolu de $[pv^2]$ serait zéro. Mais à ce minimum ne correspondrait évidemment pas la compensation la plus avantageuse. Comme nous ne connaissons pas la grandeur des erreurs v_r , nous devons

à l'égard de $[p(u-v)^2]$ chercher sa détermination la plus avantageuse, savoir sa valeur moyenne

$$\overline{[p(u-v)^2]},$$

et au minimum de cette quantité correspondra, par suite, la compensation qui doit être considérée comme la plus avantageuse.

5.

La compensation la plus avantageuse.

On a, d'après les équations (10), (11) et (12),

$$\left. \begin{aligned} O_r + v_r &= a_r[xO] + b_r[yO] + c_r[zO] + \dots, \\ O'_r &= a_r[xO'] + b_r[yO'] + c_r[zO'] + \dots, \end{aligned} \right\} \quad (21)$$

et comme $O'_r - O_r = u_r + u'_r$, on obtiendra, par soustraction de ces deux équations,

$$u_r - v_r = -u'_r + a_r[x(u + u')] + b_r[y(u + u')] + \dots \quad (22)$$

Si nous écrivons

$$* \text{ NOTE 3. } \quad [p(u-v)^2] = [pu^2] - 2[puv] + [pv^2],*$$

le dernier de ces trois termes sera déterminé par la première équation (21); au contraire, les deux premiers contiendront les erreurs inconnues. Aussi ne pouvons-nous seulement que déterminer les valeurs les plus avantageuses de ces deux termes.

On a, d'après l'équation (20),

$$[p\bar{u}^2] = nm^2, \quad (23)$$

et, d'après l'équation (22) ci-dessus,

$$\begin{aligned} p_r u_r v_r &= p_r u_r (u_r + u'_r) - p_r a_r u_r [x(u + u')] \\ &\quad - p_r b_r u_r [y(u + u')] \dots, \end{aligned}$$

par où nous trouverons facilement, au moyen de la règle de la détermination de la valeur moyenne donnée dans le paragraphe précédent,

$$[\overline{p u v}] = [\overline{p u^2}] - [\overline{p u^2} a x] - [\overline{p u^2} b y] - \dots$$

Cette expression se réduit en vertu des équations (20) à

$$n m^2 - m^2 [a x] - m^2 [b y] - \dots;$$

ensuite nous obtiendrons au moyen des équations (14)

$$[\overline{p u v}] = n m^2 - e m^2. \quad (24)$$

Par conséquent, si nous désignons la valeur moyenne cherchée de $[p(u-v)^2]$ par V , en posant

$$V = [p \overline{u^2}] - 2[p \overline{u v}] + [p v^2],$$

V sera déterminée par

$$V = [p v^2] + (2e - n)m^2. \quad (25)$$

Cela fait ressortir *qu'on doit considérer comme la plus avantageuse la compensation pour laquelle la quantité $[p v^2] + [2e - n]m^2$ et par conséquent aussi*

$$[p v^2] + 2e m^2$$

est la plus petite valeur pour une suite donnée d'observations.

Pour un nombre invariable e d'éléments, la compensation la plus avantageuse correspondra donc à la valeur minima de $[p v^2]$; c'est pourquoi nous pourrons, les coefficients a , b , c , ... étant donnés, exécuter le calcul comme à l'ordinaire; mais, si l'on compare deux développements qui sont calculés pour des nombres

différents d'éléments, on doit de plus tenir compte du nombre des éléments.

Nous pouvons aussi déterminer la valeur moyenne de $[p(u-v)^2]$ d'une autre manière, en remplaçant la vraie valeur de $[pv^2]$ par la valeur la plus avantageuse de cette somme calculée avant les observations. On a, d'après l'équation (22)

$$\begin{aligned} (\overline{u_r - v_r})^2 &= \overline{u_r^2} (1 - 2a_r x_r - 2b_r y_r \dots) \\ &\quad + a_r^2 [x^2(\overline{u^2} + \overline{u'^2})] + b_r^2 [y^2(\overline{u^2} + \overline{u'^2})] \dots \\ &\quad + 2a_r b_r [xy(\overline{u^2} + \overline{u'^2})] + \dots, \end{aligned}$$

par où l'on obtient

$$\begin{aligned} &[p(\overline{u-v})^2] \\ &= [p\overline{u^2}(1-2ax-2by\dots)] + [paa][x_2(\overline{u^2} + \overline{u'^2})] \\ &\quad + [pbb][y^2(\overline{u^2} + \overline{u'^2})] + \dots \end{aligned} \quad (26)$$

Ici l'on introduit les expressions de $x, y, z \dots$ données par (15), ce qui fait intervenir les coefficients nouveaux $\xi, \eta, \zeta \dots$ qui doivent satisfaire aux conditions (16).

Des plus, nous emploierons la notation abrégée

$$* \text{ NOTE 4. } m'^2 = \frac{1}{n-e} \left[p\overline{u^2} \left(1 - \frac{paa}{[paa]} - \frac{pbb}{[pbb]} \dots \right) \right], * \quad (27)$$

où le second membre est toujours positif, à moins qu'il n'y ait pas d'erreurs de formule, auquel cas nous aurons $m' = 0$.

L'équation (26) deviendra donc

$$\begin{aligned} [p(\overline{u-v})^2] &= (n-e)m'^2 + em^2 + [paa][\xi^2(\overline{u^2} + \overline{u'^2})] \\ &\quad + [pbb][\eta^2(\overline{u^2} + \overline{u'^2})] + \dots \end{aligned} \quad (28)$$

Cette expression aura, pour des valeurs données des coefficients $a, b, c, \dots$ sa valeur minima quand

$$\hat{\xi}_r = 0, \quad \eta_r = 0, \quad \zeta_r = 0, \quad \dots \quad (29)$$

et les coefficients $x, y, z, \dots$ seront donc, pour la détermination la plus avantageuse de (15) fournis par les formules

$$x_r = \frac{p_r a_r}{[p a a]}, \quad y_r = \frac{p_r b_r}{[p b b]}, \quad z_r = \frac{p_r c_r}{[p c c]}, \quad \dots \quad (30)$$

Si l'on emploie ces valeurs des coefficients $x, y, z, \dots$ on tirera des équations (11) et (12)

$$A = \frac{[p a O]}{[p a a]}, \quad B = \frac{[p b O]}{[p b b]}, \quad C = \frac{[p c O]}{[p c c]}, \quad \dots \quad (31)$$

et

$$[p v^2] = [p O O] - \frac{[p a O]^2}{[p a a]} - \frac{[p b O]^2}{[p b b]} - \frac{[p c O]^2}{[p c c]} \dots \quad (32)$$

Si nous voulons de nouveau introduire dans ces expressions les coefficients primitifs $a, b, c, \dots$ au lieu de $a, b, c \dots$, cela peut facilement se faire au moyen des équations (3) et (4). Si nous introduirons en même temps les constantes $\mathfrak{A}, \mathfrak{B}, \mathfrak{C}, \dots$ au lieu de $A, B, C, \dots$, ce qui pourtant sera superflu d'ordinaire dans le calcul pratique, les résultats prendront la forme sous laquelle ils sont présentés d'habitude dans la théorie des moindres carrés.

Si l'on compare les deux expressions (25) et (28) de la valeur moyenne de $[p(u-v)^2]$, on obtiendra avec les valeurs calculées des coefficients $x, y, z, \dots$, pour la détermination la plus avantageuse,

$$m^2 + m'^2 = \frac{[p v^2]}{n - e}. \quad (33)$$

L'erreur moyenne M de la détermination

$$[p u^2] - 2[p u v] = [p \bar{u}^2] - 2[p \bar{u} \bar{v}]$$

employée pour la détermination de V dans (25) est, d'après l'équation (19) rapprochée des équations (23) et (24), déterminée par

$$M^2 = ([p u^2] - 2[p u v])^2 - (2e - n)^2 m^2.$$

Le calcul de la valeur moyenne cherchée ici peut être exécuté sans difficulté et le résultat prend une forme simple, si l'on suppose que la loi exponentielle des erreurs est valable ici pour les erreurs des observations, car nous obtiendrons, grâce à ces hypothèses,

$$p_r^2 \bar{u}_r^2 = 3 m^4$$

et, comme résultat final,

$$* \text{ NOTE 5. } \quad M^2 = 2 n m^4 + 4 m^3 m^2 (n - e). * \quad (34)$$

Il en résulte

$$* \text{ NOTE 6. } \quad M \geq \sqrt{2n} \cdot m^2. *$$

Comme nous avons, de plus, d'après la détermination la plus avantageuse donnée par (33),

$$[p v^2] \geq (n - e) m^2,$$

on reconnaît que le rapport de l'erreur moyenne M , que nous employons pour la détermination de la compensation la plus avantageuse, à la somme $[p v^2]$ converge vers zéro pour un nombre croissant d'observations, si le nombre des éléments e est ou un nombre donné ou dans un rapport donné avec le nombre n des observations.

6.

La compensation la plus avantageuse calculée d'une autre manière.

La condition de la compensation la plus avantageuse trouvée dans ce qui précède ne dépend pas de la loi des erreurs de formule et le résultat ne doit donc pas varier, quelles que soient les lois de ces erreurs, à condition que nos deux suppositions relatives aux erreurs soient réalisées. Pour cette raison il ne sera pas sans intérêt de contrôler le résultat trouvé en exécutant les calculs suivant les principes appliqués dans la méthode des moindres carrés en faisant cette hypothèse, que la loi exponentielle des erreurs soit encore valable pour les erreurs de formule et que les mêmes poids, p_r , qui correspondent aux erreurs des observations, u_r , correspondent encore aux erreurs de formule, u'_r , tandis que les erreurs moyennes de l'unité de poids des deux espèces d'erreurs peuvent différer.

En partant de ces deux suppositions, on pourra dans les équations (9) et (10), savoir

$$\begin{aligned} O'_r + v'_r &= a_r A' + b_r B' + c_r C' + \dots, \\ O''_r &= a_r A'' + b_r B'' + c_r C'' + \dots, \end{aligned}$$

traiter les quantités O'_r comme observations ayant les poids p_r et O''_r comme les valeurs vraies correspondantes. Les constantes $A', B', C' \dots$ devront donc être déterminées suivant les principes connus des moindres carrés par le nombre correspondant d'„équations normales“,

$$[pav'] = 0, \quad [pbv'] = 0, \quad [pcv'] = 0 \dots \quad (35)$$

En même temps on trouvera par soustraction des deux équations ci-dessus

$$O_r + v'_r - O''_r = v'_r - u'_r = a_r(A' - A'') + b_r(B' - B'') + \dots; \quad (36)$$

par suite, au moyen des équations (35) en multipliant par $p_r v'_r$ et en additionnant toutes les équations qui correspondent aux indices 1, 2, 3 ... n , on obtiendra

$$[p v' (u' - v')] = 0. \quad (37)$$

La détermination la plus avantageuse est de plus donnée par l'équation bien connue de la théorie des moindres carrés

$$\frac{[p u'^2]}{n} = \frac{[p v'^2]}{n - e}. \quad (38)$$

Nous pourrions traiter d'une manière analogue les équations

$$\begin{aligned} O_r + v_r &= a_r A + b_r B + c_r C \dots, \\ O''_r &= a_r A'' + b_r B'' + c_r C'' \dots, \end{aligned}$$

par où nous obtiendrons

$$[p a v] = 0, \quad [p b v] = 0, \quad [p c v] = 0, \dots \quad (39)$$

et

$$O_r + v_r - O''_r = v_r - u_r - u'_r = a_r(A - A'') + b_r(B - B'') + \dots \quad (40)$$

Grâce à cette dernière équation, en multipliant par $p_r v_r$ et additionnant les équations qui correspondent aux indices 1, 2 ... n , on trouve

$$[p v (v - u - u')] = 0; \quad (41)$$

de même en multipliant par $p_r v'_r$ on trouvera d'après (35)

$$[p v' (v - u - u')] = 0, \quad (42)$$

tandis que des équations (36) et (39) on tirera d'une manière analogue

$$[p v (u' - v')] = 0. \quad (43)$$

De plus on a l'équation correspondante à (38)

$$\frac{[p(u+u')^2]}{n} = \frac{[pv^2]}{n-e}. \quad (44)$$

On peut encore ajouter deux équations que l'on obtient en remarquant que les erreurs d'observation u ne dépendent en aucune manière des erreurs u' et v' dont l'origine est due seulement au choix des coefficients $a, b, c \dots$. Il s'ensuit que la détermination la plus avantageuse et la plus probable sera donnée par les équations

$$[puu'] = 0, \quad (45)$$

$$[puv'] = 0. \quad (46)$$

Au moyen des relations, les unes exactes, les autres seulement probables, qui ont été établies entre les erreurs — et nous aurons précisément besoin de toutes ces équations — est il possible de déterminer la valeur la plus avantageuse de $[p(u-v)^2]$ exprimée par des quantités connues.

En premier lieu, en additionnant les cinq équations (37), (41), (42), (43) et (46) on trouve

$$[p(v'^2 - v^2 + uv)] = 0, \quad (47)$$

tandis que les équations (44) et (45) donneront

$$\frac{[pu^2] + [pu'^2]}{n} = \frac{[pv^2]}{n-e}. \quad (48)$$

Au moyen de ces deux équations, combinées avec l'équation (38), nous pourrons éliminer $[pu'^2]$ et $[pv'^2]$, par où nous obtiendrons

$$[puv] = \frac{n-e}{n} [pu^2].$$

La valeur la plus avantageuse de $[p(u-v)^2]$ sera donc

$$[pu^2] - 2\frac{n-e}{n}[pu^2] + [pv^2],$$

expression qui se transformera en

$$[pv^2] + (2e-n)m^2,$$

si nous introduisons la notation m , l'erreur moyenne de l'unité de poids, et si nous remplaçons $[pu^2]$ par la valeur la plus avantageuse de cette somme, à savoir nm^2 .

Le résultat est donc identique au résultat trouvé dans l'article précédent.

7.

Applications.

La condition pour qu'on puisse obtenir un résultat meilleur par les valeurs compensées que par les valeurs observées est

$$[p(u-v)^2] < [pu^2].$$

Si l'on remplace ces quantités respectivement par leurs valeurs les plus avantageuses, savoir $[pv^2] + (2e-n)m^2$ et nm^2 , on obtiendra

$$\frac{[pv^2]}{n-e} < 2m^2, \quad (49)$$

condition qui donc doit être remplie, si l'on veut qu'il y ait avantage probable à faire la compensation avec les coefficients admis.

Si l'on ne connaît pas l'erreur moyenne m de l'unité de poids des observations, mais au contraire les valeurs vraies des coefficients $a, b, c \dots$, on pourra, en faisant la compensation avec ces coefficients, déterminer la valeur la plus avantageuse de m , au moyen de l'équation

$$m^2 = \frac{[pv^2]}{n-e},$$

équation qui résulte de (33), car nous aurons dans ce cas $m' = 0$.

Au contraire, si l'on ne connaît ni les valeurs vraies des coefficients a , b , c ni l'erreur moyenne de l'unité de poids, *on ne peut pas savoir, si la compensation donnera un résultat meilleur que les observations immédiates.*

Pourtant, même dans ce cas, la compensation peut être d'une utilité essentielle, à un autre point de vue. Tandis qu'on ne peut, comme nous l'avons supposé, rien savoir de la grandeur de la somme des carrés $[pu^2]$ des erreurs d'observation elles-mêmes, on pourra au contraire déterminer par la compensation *les limites, entre lesquelles est vraisemblablement comprise la somme $[p(u-v)^2]$ des carrés des valeurs compensées* et l'on pourra de cette manière juger dans un cas donné, si l'on peut, ou non, se contenter des valeurs compensées.

On aura d'après (33)

$$m^2 = \frac{[pv^2]}{n-e} - m'^2,$$

et si nous introduisons cette valeur de m^2 dans l'expression

$$V = [pv^2] + (2e-n)m^2,$$

nous obtiendrons

$$V = \frac{e}{n-e} [pv^2] - (2e-n)m'^2, \quad (50)$$

deux expressions qui font voir que V ou la valeur la plus avantageuse de $[p(u-v)^2]$ est comprise entre les limites

$$[pv^2] \text{ et } \frac{e}{n-e} [pv^2].$$

Pour $e = \frac{1}{2}n$, ces limites coïncident.

2) *La compensation calculée comme à l'ordinaire par les valeurs vraies des coefficients $a, b, c, \dots$ est plus avantageuse, qu'aucune autre compensation à un nombre d'éléments égal ou plus grand.* Par conséquent, si l'on n'a qu'un élément, les valeurs vraies des coefficients donneront toujours la compensation la plus avantageuse.

On reconnaît facilement l'exactitude de cette proposition, en posant dans l'une des expressions ci-dessus de V

$$[pv^2] = (n-e)(m^2 + m'^2),$$

par où l'on obtient

$$V = em^2 + (n-e)m'^2. \quad (51)$$

Si l'on se sert dans la compensation des valeurs vraies des coefficients $a, b, c, \dots$, on aura $m' = 0$ et par conséquent

$$V = em^2;$$

au contraire si l'on se sert d'autres valeurs et si l'on admet e' éléments, on aura

$$V = e'm^2 + (n-e)m'^2,$$

valeur qui est plus grande que la précédente, em^2 , si e' est égal ou supérieur à e .

Au contraire, on peut obtenir une compensation plus avantageuse avec un nombre plus petit d'éléments.

D'après (32), nous avons

$$[pv^2] - [pOO] = -\frac{[paO]^2}{[paa]} - \frac{[pbO]^2}{[pbb]} - \dots - \frac{[plO]^2}{[pll]},$$

en désignant par l le dernier des éléments $a, b, c \dots$ tandis que leur nombre est désigné comme dans ce qui précède par e .

Si nous négligeons le dernier terme de cette série, elle sera augmentée de $\frac{[p l O]^2}{[p l l]}$, et le nombre des éléments diminuera d'une unité. Par conséquent, en négligeant le dernier élément, on augmentera V de

$$\frac{[p l O]^2}{[p l l]} - 2m^2.$$

Donc, si l'on a

$$\frac{[p l O]^2}{[p l l]} < 2m^2$$

cet accroissement sera négatif, et l'on peut vraisemblablement avec avantage négliger ce dernier élément. Que les observations puissent en réalité avoir de telles valeurs, c'est ce qui se voit facilement; rien même n'empêche d'imaginer une série observations pour laquelle $[p l O] = 0$.

Il va sans dire qu'on peut de même se demander, si l'on ne doit pas négliger deux ou plusieurs éléments et remplacer les coefficients $a, b, c \dots$ que l'on conserve par des valeurs autres que celles qui originairement étaient les valeurs vraies.

3) Si l'on ne connaît pas les valeurs vraies des coefficients $a, b, c \dots$, on doit choisir ces valeurs arbitrairement, de manière pourtant à satisfaire à la connexion donnée entre elles, d'où dépend la compensation (voir section 1).

Ensuite, on calcule successivement les coefficients $a, b, c \dots$ (voir section 2), et l'on trouve de cette manière les termes successifs de la série $[p v^2] - [p O O]$ indiquée ci-dessus.

Quand on parvient à un terme numériquement plus petit que $2m^2$, il doit être négligé et de même l'élément correspondant; en général on pourra s'en tenir aux termes calculés avant celui-là.

Si, de plus, m est inconnu, on pourra pourtant en tout cas reconnaître quels termes de la série *on ne doit pas négliger*: ainsi tous les termes qui sont numériquement plus grands que $2m^2$, et par conséquent aussi les termes qui sont plus grands que $2 \frac{[pv^2]}{n-e}$ doivent être conservés. On pourra ici déterminer $[pv^2]$ au moyen des termes déjà calculés. En général il ne sera pas de grande conséquence pour la détermination du nombre des éléments, que l'on connaisse m , ou non.

Finalement, si l'on veut comparer plusieurs compensations calculées avec des coefficients différents, on doit en premier lieu déterminer pour chaque compensation en particulier le nombre le plus avantageux d'éléments et puis comparer les valeurs de V obtenues de cette manière par les différentes compensations.

Parfois on peut, en outre, demander s'il est plus avantageux de calculer toute la série des observations avec les mêmes coefficients ou de la subdiviser en plusieurs séries et de calculer chaque partie au moyen de coefficients particuliers, auquel cas on doit remarquer que, les nombres des éléments des séries particulières étant e' , e'' ..., le nombre total pour toute la série sera $e' + e'' + \dots$.

Une telle subdivision d'une série donnée d'observations peut, par exemple, être encore appliquée à des observations de deux ou de plusieurs quantités, dans lesquelles les valeurs observées sont mêlées d'une manière

inconnue, ou à des observations d'une seule quantité, pour lesquelles l'on connaît l'erreur moyenne des erreurs accidentelles, mais où se sont en outre introduites des erreurs constantes pour quelques observations particulières, sans qu'on sache pour lesquelles. Nous ne traiterons pas ces problèmes particuliers de compensation, parce que nous avons indiqué la marche à suivre pour trouver une solution, s'il en existe une, le problème étant en tout cas réduit à la détermination des corrections ϵ de manière à obtenir pour la série totale d'observations la valeur la plus petite possible de

$$[pv^2] + 2em^2.$$

NOTES.

NOTE 1. Ce mémoire contient un supplément à la théorie ordinaire de la compensation des erreurs d'observation. En se mettant au point de vue aprioristique, Lorenz soutient que, même dans le cas où l'on connaît les valeurs vraies des éléments, on n'obtiendra pas toujours la compensation la plus avantageuse en se servant de ces valeurs vraies.

La théorie de Lorenz a été combattue par M. le capitaine (à présent général) Zachariae.*

NOTE 2. La condition

$$[p(u-v)^2] = \text{Min.}$$

donnera le même résultat que la théorie ordinaire des moindres carrés, si l'on remplace, comme le fait Lorenz, les quantités inconnues par leurs valeurs moyennes.

NOTE 3. Le texte de Lorenz en ce passage est très peu clair. Lorenz dit que la valeur de $[pv^2]$ est déterminée par l'équation (21) et que, pour cette raison, on ne peut déterminer la valeur la plus avantageuse de

* Voir „Tidsskrift for Mathematikk“. 1872. Zachariae: Udjevning af Iagttagelsesfejl (deux notes) p. 97—103 et 182—188, et les répliques de Lorenz, dans le même journal, Bidrag til Udjevningens problemet, 1872, p. 125—135 et Et Gjensvar, 1873, p. 31—32.

$[pu^2] - 2[puv] + [pv^2]$ que par les deux premiers termes. On doit donc en conclure qu'il considère v_r comme une constante et que par conséquent $\overline{[puv]} = 0$. Le résultat ne concorde pas avec celui de Lorenz; mais c'est qu'en effet, ici comme dans la formule (28), il remplace la vraie valeur de v_r par sa valeur la plus avantageuse calculée avant l'observation.

NOTE 4. Je ne peux pas voir pourquoi la quantité désignée par Lorenz par m'^2 doit être supposée positive.

NOTE 5. L'évaluation de l'expression

$$M^2 = (\overline{[pu^2]} - 2\overline{[puv]})^2 - (2e - n)^2 m^4$$

peut se faire de la manière suivante.

On a

$$M^2 = \overline{[pu^2]^2} + 4\overline{[puv][pu(v-u)]} - (2e - n)^2 m^4;$$

nous allons calculer séparément la valeur de chaque terme.

On a

$$[pu^2] = p_1 u_1^2 + p_2 u_2^2 + \dots + p_n u_n^2$$

et par conséquent

$$[pu^2]^2 = [p^3 u^4] + 2 \sum p_r p_s u_r^2 u_s^2,$$

où le dernier groupe comprend $\frac{n(n-1)}{2}$ termes, et comme $\overline{p_r^2 u^4} = 3m^4$, $\overline{p_r u^2} = m^2$, on aura

$$\overline{[pu^2]^2} = n(n+2)m^4.$$

Considérons ensuite le terme

$$\begin{aligned} & \overline{[puv][pu(v-u)]} \\ &= \overline{[pu(u+u') - a[x(u+u')] - b[y(u+u')] \dots]} \\ & \times \overline{[pu(u'-a[x(u+u')] - b[y(u+u')] \dots]}. \end{aligned}$$

Tout terme qui ne contient u' qu'à la première puissance s'évanouira; aussi les termes qui dépendent de u' se trouvent ils dans l'expression

$$\begin{aligned} & \overline{[pu(u' - a[xu'] - b[yu']) \dots]^2} \\ = & \overline{[p^2u^2(u'^2 + a^2[x^2u'^2] + b^2[y^2u'^2] - 2axu'^2 - 2byu'^2 \dots)]} \\ = & m^2m'^2(n - e). \end{aligned}$$

Reste à calculer la valeur de

$$\begin{aligned} & - \overline{[pu(u - a[xu] - b[yu]) \dots] [pu(a[xu] + b[yu]) \dots]} \\ = & - \overline{[pu^2] [paxu^2 + pbyu^2 \dots]} + \overline{[pu(a[xu] + b[yu]) \dots]^2}. \end{aligned}$$

Mais on peut remplacer $p_r u_r^2$ par m^2 , $p_r^2 u_r^4$ par $3m^4$; par conséquent, on aura

$$\begin{aligned} & \overline{[pu^2] [paxu^2 + pbyu^2 \dots]} \\ = & m^4(n + 2)([ax] + [by] + \dots), \end{aligned}$$

et comme $[ax] = [by] = \dots 1$, il viendra

$$\overline{[pu^2] [paxu^2 + pbyu^2 \dots]} = m^4 e(n + 2).$$

Finalement on aura

$$\begin{aligned} & \overline{[pu(a[xu] + b[yu]) \dots]^2} \\ = & \overline{[pu a]^2 [xu]^2} + \overline{[pu b]^2 [yu]^2} + \dots \\ & + 2 \overline{[pu a] [pu b] [xu] [yu]} + \dots \\ = & m^4 \left([a^2 p] \left[\frac{x^2}{p} \right] + [b^2 p] \left[\frac{y^2}{p} \right] + \dots \right. \\ & + 4 \sum (a_r a_s x_r x_s + b_r b_s y_r y_s + \dots) \\ & + 2([a^2 x^2] + [b^2 y^2] + \dots) \\ & + 2 \sum [p a b] \left[\frac{xy}{p} \right] \\ & + 4 \sum [xy b a] \\ & \left. + 2 \sum \sum a_r b_s x_r y_s + 2 \sum \sum a_r b_s x_s y_r \right) \\ = & m^4 \left(\sum [a^2 p] \left[\frac{x^2}{p} \right] + 2 \sum [a x]^2 + 2 \sum [a y] [b x] \right. \\ & \left. + 2 \sum [a x] [b y] + 2 \sum [p a b] \left[\frac{xy}{p} \right] \right). \end{aligned}$$

Mais comme $[pab] = [pac] = \dots = 0$ et, par conséquent, $[ay] = [az] = \dots = 0$, on verra facilement que toute l'expression se réduit à

$$m^4(3e + e(e-1)).$$

On aura donc

$$\begin{aligned} M^2 &= n(n+2)m^4 + 4m^2m'^2(n-e) - 4m^4e(n+2) \\ &\quad + 4m^4(3e + e(e-1)) - (2e-n)^2m^4 \\ &= 2nm^4 + 4m^2m'^2(n-e). \end{aligned}$$

NOTE 6. Lorenz a, par méprise, écrit

$$M \leq \sqrt{2n} \cdot m^2;$$

mais sa conclusion est pourtant juste.



SUR LA RÉDUCTION
DU FACTEUR EULÉRIEN.

SUR LA RÉDUCTION DU FACTEUR EULÉRIEN.* * NOTE 1.

TIDSSKRIFT FOR MATEMATIK 1874, P. 33—39.

Les fonctions transcendantes qui seront traitées ici peuvent être classifiées de la manière suivante. La transcendante monôme du premier ordre est définie par

$$\theta_1 = \int f(\alpha) d\alpha,$$

où la limite supérieure de l'intégrale est une fonction algébrique de x et y , tandis que la limite inférieure est une constante. La fonction f est, comme les autres fonctions qui, dans ce qui suit, seront désignées par la même lettre affectée d'indices différents, une fonction algébrique de la variable indépendante, ici α ; au contraire, on suppose que θ ne puisse pas être exprimée algébriquement par la variable.

Une fonction transcendante du premier ordre qui est désignée par t_1 est une fonction algébrique de x , y , transcendantes monômes du premier ordre comme θ_1 , $\theta'_1 \dots$ et des fonctions inverses de θ_1 , $\theta'_1 \dots \eta_1$, $\eta'_1 \dots$ considérées comme fonctions de α , $\alpha_1 \dots$. *Une fonction monôme du second ordre a la forme*

$$\theta_2 = \int f_1(t_1) dt_1.$$

La fonction transcendante du second ordre est une fonction algébrique de x , y , θ_1 , $\theta'_1 \dots \eta_1$, $\eta'_1 \dots$, θ_2 , $\theta'_2 \dots \eta_2$, $\eta'_2 \dots$, où les dernières fonctions sont les fonctions

*NOTE 2. inverses de $\theta_2, \theta'_2 \dots$ * En continuant de cette manière, la définition des transcendentes monômes et des fonctions transcendantes d'ordre supérieur est immédiatement évidente.

Une équation entre les fonctions transcendantes de ladite espèce écrite sous la forme la plus simple, de manière que le nombre des transcendentes monômes de l'ordre le plus élevé et de leurs fonctions inverses soit réduit à son minimum, doit être identique en particulier par rapport à chaque fonction de ladite espèce; car autrement on pourrait exprimer une d'elles algébriquement par les autres et par les fonctions d'ordre inférieur, de sorte que le nombre des transcendentes de l'ordre le plus élevé ne serait pas réduit à son minimum.

Je suppose maintenant qu'une équation donnée $dy + Pdx = 0$, où P est une fonction algébrique de x et y , a un facteur intégrant qui peut être exprimé par les transcendentes de l'espèce définie ci-dessus; le problème à résoudre sera, de trouver une *forme normale* à laquelle peut être réduit chaque facteur transcendant de ladite espèce. J'ai cherché à résoudre ce problème de cette manière: j'imagine qu'un facteur intégrant F soit donné comme fonction transcendante de l'ordre n , c'est à dire comme fonction algébrique de $\theta_n, \theta'_n, \dots, \eta_n, \eta'_n, \dots$, de fonctions monômes d'ordre inférieur, de leurs fonctions inverses et enfin de x et y . Je cherche alors à réduire ce facteur à un autre, d'ordre inférieur et ainsi de suite. On suppose toujours dans le calcul que le nombre des transcendentes est réduit à son minimum.

Nous chercherons d'abord la réduction de F par rapport à θ_n , que nous désignerons provisoirement, pour abrégé, par la seule lettre θ .

L'équation de condition du facteur est

$$\frac{\partial F}{\partial x} + P \frac{\partial F}{\partial y} + F \frac{\partial P}{\partial y} = 0. \quad (1)$$

Comme F est une fonction algébrique de θ (et des autres variables), nous pourrions considérer F comme racine d'une équation algébrique irréductible

$$a_0 + a_1 F + \dots + a_n F^n = 0. \quad (2)$$

où les coefficients sont des fonctions entières et rationnelles de θ (et des autres variables). De la même manière P est racine d'une équation algébrique; mais à une racine particulière de celle-ci ne correspond pas nécessairement une racine quelconque de (2), et il n'est permis que de supposer qu'une seule racine de (2) satisfasse à l'équation différentielle (1).

Soit $Q\theta^s$ la valeur limite vers laquelle tend une racine de l'équation (2), lorsque θ croît à l'infini ou, ce qui revient au même, le premier terme de la série qu'on obtient par le développement de la racine suivant les puissances décroissantes de θ .

Ce terme peut être déterminé de la manière suivante. Dans l'équation (2), F est remplacée par $Q\theta^s$ et l'on sépare dans les coefficients $a_0, a_1, \dots$ les termes qui contiennent les puissances les plus élevées de θ , à savoir $c_0\theta^{\mu_0}, c_1\theta^{\mu_1}, \dots$; nous aurons ainsi provisoirement l'équation suivante pour déterminer Q

$$c_0\theta^{\mu_0} + c_1\theta^{\mu_1+s}Q + \dots + c_m\theta^{\mu_m+ms}Q^m = 0. \quad (3)$$

Pour* trouver les valeurs différentes que peut prendre *NOTE 3.
 Q , nous chercherons d'abord l'exposant μ_q pour lequel

$$\frac{\mu_q - \mu_0}{q} \geq \frac{\mu_p - \mu_0}{p},$$

où p désigne un indice quelconque plus grand que 0. Puis on peut poser $s = -\frac{\mu_q - \mu_0}{q}$, valeur pour laquelle $\mu_p + ps \leq \mu_0$, par où l'on reconnaît qu'aucun exposant de θ dans (3) n'est plus grand que μ_0 . En négligeant tous les termes dont les exposants sont plus petits que μ_0 , θ^{μ_0} sera facteur commun des termes restants et le coefficient de θ^{μ_0} égalé à zéro donnera une équation propre à déterminer Q .

Si q est l'indice le plus grand des termes de cette dernière équation, on cherchera ensuite un exposant μ_r , à indice plus grand, pour lequel

$$\frac{\mu_r - \mu_q}{r - q} > \frac{\mu_p - \mu_q}{p - q},$$

où p désigne un indice quelconque plus grand que q . Puis on pose $s = -\frac{\mu_r - \mu_q}{r - q}$, par où l'on obtiendra

$$\mu_q + qs \geq \mu_p + ps,$$

et par conséquent $\mu_q + qs$ sera à présent la puissance la plus grande de θ qui entre dans (3). En négligeant tous les termes de degré inférieur et en divisant par $\theta^{\mu_q + qs} Q^q$, on obtiendra une équation nouvelle par laquelle d'autres valeurs de Q peuvent être déterminées. En continuant de cette manière, on obtiendra les équations suffisantes pour la détermination de toutes les valeurs de Q correspondant aux m racines de l'équation (2).

Du reste ces équations peuvent aisément être trouvées de la manière suivante. Qu'on représente $\mu_0, \mu_1 \dots$ comme des ordonnées successives rangées d'après leurs indices et que l'on construise un polygone convexe, de

manière à faire coïncider ses sommets avec les extrémités des ordonnées, de telle manière qu'aucune d'elles ne soit en dehors du polygone. A chaque côté du polygone passant par deux ou plusieurs points μ correspondra alors une équation en Q ; par exemple, à un côté passant par les points μ_2, μ_4, μ_5 correspondra l'équation

$$c_2 + c_4 Q^2 + c_5 Q^3 = 0.$$

Après avoir trouvé de cette manière la valeur de $Q\theta^s$ correspondante à F dans l'équation (1), nous remplaçons F par cette quantité, et, comme l'équation est identique par rapport à θ , nous égalons à zéro le coefficient de la plus haute puissance de θ . Comme $\frac{\partial \theta}{\partial x}$ et $\frac{\partial \theta}{\partial y}$ ne contiennent que des transcendentes d'ordre inférieur, cette puissance est θ^s , et l'on trouvera

$$\frac{\partial Q}{\partial x} + P \frac{\partial Q}{\partial y} + \frac{\partial P}{\partial y} Q = 0,$$

par où l'on reconnaît que Q , qui ne contient pas θ , est un facteur intégrant.

De la même manière on peut réduire le facteur par rapport à toutes les transcendentes monômes d'ordre n ; et l'on obtiendra finalement un facteur ne contenant que les fonctions inverses des transcendentes monômes d'ordre n .

Si une telle fonction est désignée par η , et si l'on se sert de nouveau de la lettre F pour représenter le facteur intégrant, l'équation de condition sera

$$\frac{\partial F}{\partial x} + P \frac{\partial F}{\partial y} + \frac{1}{f(\eta)} \frac{\partial F}{\partial \eta} \left[\frac{\partial t_{n-1}}{\partial x} + P \frac{\partial t_{n-1}}{\partial y} \right] + \frac{\partial P}{\partial y} F = 0, \quad (4)$$

où la différentiation par rapport à η est indiquée sé-

parément et où η est la fonction inverse de $\theta = \int f(t_{n-1}) dt_{n-1}$, d'où résulte

$$\frac{\partial \eta}{\partial t_{n-1}} = \frac{1}{f(\eta)}.$$

Dans l'équation ci-dessus nous chercherons le coefficient de la plus haute puissance de η , après avoir déterminé de la manière indiquée ci-dessus le premier terme de $Q\eta^s$ dans le développement de F suivant les puissances décroissantes de η . En même temps on trouvera la quantité $k\eta^\sigma$ où k est une constante, savoir le premier terme du développement correspondant de $\frac{1}{f(\eta)}$. Les termes de l'équation (4) qui entreront en considération seront donc

$$\left. \begin{aligned} \eta^s \left[\frac{\partial Q}{\partial x} + P \frac{\partial Q}{\partial y} + sk Q \eta^{\sigma-1} \left(\frac{\partial t_{n-1}}{\partial x} + P \frac{\partial t_{n-1}}{\partial y} \right) \right] \\ + \frac{\partial P}{\partial y} Q \Big] + \dots = 0. \end{aligned} \right\} \quad (5)$$

Il s'ensuit que Q sera un facteur intégrant, si σ est plus petit que 1, ou si s est égal à zéro. Si $\sigma > 1$ et si s diffère de zéro, on aura

$$\frac{\partial t_{n-1}}{\partial x} + P \frac{\partial t_{n-1}}{\partial y} = 0, \quad (6)$$

par où l'on voit, en différentiant l'équation par rapport à y , que $\frac{\partial t_{n-1}}{\partial y}$ doit être elle-même un facteur intégrant.

Enfin, si $\sigma = 1$, on doit tenir compte de tous les termes de l'équation (5), et l'équation qu'on obtient en égalant le coefficient de η^s à zéro fait voir que $Q\eta^{skt_{n-1}}$ sera un facteur intégrant; car, si l'on remplace F par cette valeur dans l'équation de condition (1), on obtiendra l'équation trouvée.

Après réduction par rapport à toutes les transcendentes inverses d'ordre n , le facteur intégrant peut être mis sous la forme $\varphi_{n-1} e^{\psi_{n-1}}$ ou $e^{\psi_{n-1} + l\varphi_{n-1}}$, où φ_{n-1} et ψ_{n-1} sont des fonctions transcendentes d'ordre $(n-1)$. Néanmoins, en vue des calculs qui suivront, ce facteur sera rapporté à la forme plus générale

$$F = e^{\psi_{n-1}} + C_n \theta_n + C'_n \theta'_n + \dots, \quad (7)$$

où $C_n, C'_n \dots$ sont des constantes.

L'exposant de e , que nous désignerons pour abréger par φ , a pour dérivée par rapport à une variable quelconque une fonction transcendente d'ordre inférieur à n . Dans l'équation de condition de φ

$$\frac{\partial \varphi}{\partial x} + P \frac{\partial \varphi}{\partial y} + \frac{\partial P}{\partial y} = 0, \quad (8)$$

les transcendentes monômes et leurs fonctions inverses seront au plus d'ordre $n-1$. Par conséquent, cette équation est identique par rapport à θ_{n-1} et peut donc être différenciée par rapport à cette variable. Comme les coefficients différentiels de θ_{n-1} , de même que P , ne contiennent pas θ_{n-1} , l'équation deviendra après avoir été différenciée*

* NOTE 4.

$$\frac{\partial \left(\frac{\partial \varphi}{\partial \theta_{n-1}} \right)}{\partial x} + P \frac{\partial \left(\frac{\partial \varphi}{\partial \theta_{n-1}} \right)}{\partial y} = 0. \quad (9)$$

Par suite, si $\frac{\partial \varphi}{\partial \theta_{n-1}}$ n'est pas une constante*, $\frac{\partial \left(\frac{\partial \varphi}{\partial \theta_{n-1}} \right)}{\partial y}$ * NOTE 5.

sera un facteur intégrant; celui-ci est donc réduit à l'ordre $n-1$. De même on peut faire une réduction analogue, si $\frac{\partial \varphi}{\partial \theta'_{n-1}}$ n'est pas constant etc.; ce n'est que dans le cas,

où φ a la forme

$$\varphi = \psi_{n-1} + C_{n-1} \theta_{n-1} + C'_{n-1} \theta'_{n-1} + \dots \quad (10)$$

ψ_{n-1} ne contenant pas θ_{n-1} , θ'_{n-1} ... mais seulement leurs fonctions inverses, qu'on ne peut plus faire aucune réduction ultérieure.

Si l'on différentie l'équation (8) par rapport à η_{n-1} , que nous désignerons pour abréger par la seule lettre η ,
NOTE 6. et si l'on divise par $f(\eta)$, on trouvera*

$$\frac{\delta \left(\frac{\partial \varphi}{\partial \eta} \cdot \frac{1}{f(\eta)} \right)}{\partial x} + P \frac{\delta \left(\frac{\partial \varphi}{\partial \eta} \cdot \frac{1}{f(\eta)} \right)}{\partial y} = 0. \quad (11)$$

Par cette équation on peut de nouveau réduire le facteur intégrant à l'ordre $n-1$, à moins que l'on n'ait $\frac{\partial \varphi}{\partial \eta} \cdot \frac{1}{f(\eta)} = c$, d'où l'on déduira $\varphi = \psi + c \int f(\eta) d\eta$, ψ étant indépendant de η . Mais on reconnaît facilement, que la dernière intégrale est d'ordre moindre que η , de manière que φ ne contient pas en réalité η dans ce cas.

Nous pouvons donc conclure, ou bien que la fonction ψ_{n-1} de l'équation (10) ne doit contenir aucune des transcendentes inverses d'ordre $(n-1)$, et φ aura alors la forme

$$\varphi = \psi_{n-2} + C_{n-1} \theta_{n-1} + C'_{n-1} \theta'_{n-1} \dots, \quad (12)$$

ou bien que la facteur intégrant peut être réduit à l'ordre $(n-1)$. Mais on pourra alors continuer la réduction de la même manière que ci-dessus jusqu'à ce que nous arrivions à la forme de φ indiquée dans (12).

* C'est M. le Dr. phil. Jul. Petersen qui a le premier remarqué cette équation.

Mais on reconnaît alors facilement que la réduction peut être continuée jusqu'à ce que nous arrivions finalement à la forme

$$\varphi = \phi + C_1 \theta_1 + C'_1 \theta'_1 \dots$$

où ϕ est une fonction algébrique de x et y .

Donc, si le facteur Eulérien de l'équation $dy + Pdx = 0$ peut être exprimé par des transcendentes de l'espèce considérée ici, ce facteur peut être réduit à une forme normale déterminée par

$$F = e^{\psi + \int f(\alpha) d\alpha + \int f'(\alpha') d\alpha' + \dots}$$

Les sens des notations employées dans cette équation est mis en évidence par ce qui précède.

NOTES.

NOTE 1. Ce mémoire est un remaniement d'une note „Om den Eulerske Factor“ (sur le facteur Eulérien) insérée dans Math. Tidss., 1873, p. 174—176.

Comme M. le professeur J. Petersen avait averti Lorenz d'une erreur contenue dans ladite note, tandis que le résultat conservait sa validité, Lorenz a écrit le présent mémoire.

NOTE 2. Les fonctions inverses η_n sont définies de la manière suivante.

Soit $\theta_n = \int f(t_{n-1}) dt_{n-1}$, où t_{n-1} est une fonction transcendante d'ordre $n-1$; alors on aura

$$t_{n-1} = \int f(\eta_n) d\eta_n,$$

ou bien

$$\frac{\partial \eta_n}{\partial t_{n-1}} = \frac{1}{f(\eta_n)}.$$

NOTE 3. Lorenz cherche les valeurs de s pour lesquelles deux termes au moins ont des exposants égaux et supérieurs à tous les autres.

Le plus grand exposant doit au moins être égal à μ_0 , et l'on cherche à déterminer une valeur de s pour laquelle $\mu_0 = \mu_q + sq$ est supérieur ou au moins égal aux valeurs des autres exposants. On doit donc avoir

$$\mu_0 = \mu_q + sq \geq \mu_p + sp,$$

où p peut être un nombre quelconque compris entre 0 et m . Par suite on aura

$$s = \frac{\mu_0 - \mu_q}{q} \leq \frac{\mu_0 - \mu_p}{p}, \quad (a)$$

et si cette condition est remplie, s peut avoir la valeur indiquée par Lorenz $\frac{\mu_0 - \mu_q}{q}$.

$\frac{\mu_0 - \mu_q}{q}$ est la valeur minimum de s . Supposons que q soit la plus grande valeur pour laquelle $\mu_0 = \mu_q + qs$; alors on aura,

$$\mu_p + ps < \mu_q + qs,$$

si $s > \frac{\mu_0 - \mu_q}{q}$ et $p < q$, et l'on obtiendra une nouvelle valeur de s en cherchant un nombre $r > q$, pour lequel

$$\mu_q + sq = \mu_r + sr \geq \mu_p + sp, \quad (b)$$

p étant un nombre quelconque compris entre les limites q et m . Si cette condition est remplie, on peut poser

$$s = \frac{\mu_q - \mu_r}{r - q} \leq \frac{\mu_q - \mu_p}{p - q}.$$

On reconnaît facilement qu'on peut continuer de la même manière en se servant à présent de μ_r comme point de départ, si r est le plus grand nombre pour lequel la condition (b) est remplie, et ainsi de suite.

Le raisonnement de Lorenz n'est pourtant pas exact, car on ne sait pas si les coefficients a ne contiennent pas de variables qui deviennent infinies en même temps que θ .

NOTE 4. On voit immédiatement que

$$\frac{\partial \left(\frac{\partial \varphi}{\partial x} \right)}{\partial \theta_{n-1}} + P \frac{\partial \left(\frac{\partial \varphi}{\partial y} \right)}{\partial \theta_{n-1}} = 0;$$

mais on peut écrire

$$\frac{\partial \varphi}{\partial x} = \frac{\partial \varphi}{\partial \theta_{n-1}} \frac{\partial \theta_{n-1}}{\partial x} + \left[\frac{\partial \varphi}{\partial x} \right],$$

où $\left[\frac{\partial \varphi}{\partial x} \right]$ désigne la dérivée de φ par rapport à x , θ_{n-1} étant considérée comme une variable indépendante de x . Puis on aura

$$\frac{\partial \left(\frac{\partial \varphi}{\partial x} \right)}{\partial \theta_{n-1}} = \frac{\partial^2 \varphi}{\partial \theta_{n-1}^2} \frac{\partial \theta_{n-1}}{\partial x} + \left[\frac{\partial \left(\frac{\partial \varphi}{\partial \theta_{n-1}} \right)}{\partial x} \right] = \frac{\partial \left(\frac{\partial \varphi}{\partial \theta_{n-1}} \right)}{\partial x},$$

et de même on obtiendra la valeur de $\frac{\partial \left(\frac{\partial \varphi}{\partial y} \right)}{\partial \theta_{n-1}}$.

NOTE 5. $\frac{\partial \left(\frac{\partial \varphi}{\partial \theta_{n-1}} \right)}{\partial y}$ sera un facteur intégrant si $\frac{\partial \varphi}{\partial \theta_{n-1}}$ n'est pas indépendant de y , c'est à dire ne se réduit pas à une constante ou à une fonction de x seul. Lorenz n'a pas considéré le dernier cas.

NOTE 6. Si η_{n-1} est définie par l'équation

$$t_{n-2} = \int f(\eta_{n-1}) d\eta_{n-1},$$

on aura

$$\frac{\partial \eta_{n-1}}{\partial t_{n-2}} = \frac{1}{f(\eta_{n-1})},$$

et, par conséquent,

$$\frac{\partial \varphi}{\partial x} = \frac{\partial \varphi}{\partial \eta_{n-1}} \frac{1}{f(\eta_{n-1})} \cdot \frac{\partial t_{n-2}}{\partial x} + \left[\frac{\partial \varphi}{\partial x} \right],$$

où $\left[\frac{\partial \varphi}{\partial x}\right]$ désigne la dérivée de φ par rapport à x , prise en regardant η_{n-1} comme une variable indépendante. Puis on aura

$$\frac{\partial \left(\frac{\partial \varphi}{\partial x}\right)}{\partial \eta_{n-1}} = \frac{\partial \left(\frac{\partial \varphi}{\partial \eta_{n-1}} \frac{1}{f(\eta_{n-1})}\right)}{\partial \eta_{n-1}} \cdot \frac{\partial t_{n-2}}{\partial x} + \frac{\partial \left[\frac{\partial \varphi}{\partial x}\right]}{\partial \eta_{n-1}},$$

et, après multiplication par $\frac{1}{f(\eta_{n-1})}$,

$$\begin{aligned} & \frac{1}{f(\eta_{n-1})} \frac{\partial \left(\frac{\partial \varphi}{\partial x}\right)}{\partial \eta_{n-1}} \\ &= \frac{1}{f(\eta_{n-1})} \frac{\partial \left(\frac{\partial \varphi}{\partial \eta_{n-1}} \frac{1}{f(\eta_{n-1})}\right)}{\partial \eta_{n-1}} \frac{\partial t_{n-2}}{\partial x} + \frac{1}{f(\eta_{n-1})} \frac{\partial \left[\frac{\partial \varphi}{\partial x}\right]}{\partial \eta_{n-1}} \\ &= \frac{\partial \left(\frac{1}{f(\eta_{n-1})} \frac{\partial \varphi}{\partial \eta_{n-1}}\right)}{\partial \eta_{n-1}} \frac{\partial \eta_{n-1}}{\partial x} + \frac{\partial \left[\frac{1}{f(\eta_{n-1})} \frac{\partial \varphi}{\partial \eta_{n-1}}\right]}{\partial x} \\ &= \frac{\partial \left(\frac{1}{f(\eta_{n-1})} \frac{\partial \varphi}{\partial \eta_{n-1}}\right)}{\partial x}. \end{aligned}$$

De la même manière on obtiendra

$$\frac{1}{f(\eta_{n-1})} \frac{\partial \left(\frac{\partial \varphi}{\partial y}\right)}{\partial \eta_{n-1}} = \frac{\partial \left(\frac{1}{f(\eta_{n-1})} \frac{\partial \varphi}{\partial \eta_{n-1}}\right)}{\partial y}.$$

ÉQUATIONS
CINÉTIQUES FONDAMENTALES
D'UN SYSTÈME DE POINTS.

EQUATIONS CINÉTIQUES FONDAMENTALES D'UN SYSTÈME DE POINTS.

TIDSSKRIFT FOR MATEMATIK 1875, P. 81—86.

Un point donné, fixe dans l'espace, étant déterminé par ses coordonnées a, b, c , convenons que le nombre des points appartenant à un système donné qui se trouvent, à l'instant t , à l'intérieur du parallépipède limité par les plans a et $a + da$, b et $b + db$, c et $c + dc$ soit représenté par

$$N(t, a, b, c) d\varpi.$$

Les points peuvent avoir des vitesses différentes. Ceux d'entre eux dont les vitesses composantes sont comprises entre les limites a' et $a' + da'$, b' et $b' + db'$, c' et $c' + dc'$ seront supposés au nombre de

$$N(t, a, b, c) s(t, a, b, c, a', b', c') d\varpi d\varpi',$$

où $d\varpi' = da' db' dc'$. La fonction s peut être appelée, d'un terme emprunté au calcul des probabilités, la *probabilité* des composantes de vitesse a', b', c' dans le parallépipède ϖ et au moment t . D'une manière analogue, soit $s_1(t, a, b, c, a', b', c', a'', b'', c'')$ la probabilité des composantes d'accélération a'', b'', c'' qui correspondent aux composantes de vitesse a', b', c' , aux coordonnées de l'espace a, b, c et au moment t , et ainsi de suite.

Lorsqu'un point d'un système qui, au moment donné t , a pour coordonnées x, y, z et pour vitesses composantes x', y', z' , satisfait aux conditions

$$a < x + x'\tau + \dots < a + da, \quad b < y + y'\tau + \dots < b + db, \\ c < z + z'\tau + \dots < c + dc,$$

il se trouvera, au bout du temps $t + \tau$, à l'intérieur du parallélépipède ϖ . En négligeant les puissances supérieures de τ , nous pourrons déterminer le nombre total des points qui au moment $t + \tau$ se trouveront à l'intérieur du parallélépipède $d\varpi$ par l'intégrale

$$N(t + \tau, a, b, c) d\varpi \\ = \iiint dx' dy' dz' \iiint dx dy dz N(t, x, y, z) s(t, x, y, z, x', y', z'),$$

où les intégrations par rapport à x, y, z doivent être effectuées entre les limites $a - x'\tau$ et $a + da - x'\tau$, $b - y'\tau$ et $b + db - y'\tau$, $c - z'\tau$ et $c + dc - z'\tau$; celles qui se rapportent à x', y', z' entre les limites à l'extérieur desquelles s s'évanouit, car ces variables doivent parcourir toutes les valeurs possibles.

De cette manière on obtiendra d'abord

$$N(t + \tau, a, b, c) = \iiint dx' dy' dz' N(t, x, y, z) s(t, x, y, z, x', y', z')$$

où

$$x = a - x'\tau, \quad y = b - y'\tau, \quad z = c - z'\tau,$$

et puis, en comparant les coefficients de τ ,

$$\frac{\partial N}{\partial t} + \iiint dx' dy' dz' \left[\frac{\partial (Nsx')}{\partial a} + \frac{\partial (Nsy')}{\partial b} + \frac{\partial (Nsz')}{\partial c} \right] = 0.$$

Pour abrégér, on écrit ici N au lieu de $N(t, a, b, c)$ et s au lieu de $s(t, a, b, c, x' y' z')$.

Nous appellerons l'intégrale

$$\iiint dx' dy' dz' sx'$$

la valeur moyenne de x' correspondante au point a, b, c , et nous la désignerons par $\overline{a'}$, et nous désignerons de même les valeurs moyennes de b' et c' par $\overline{b'}$ et $\overline{c'}$. Par

l'emploi de ces notations, l'équation ci-dessus prendra la forme

$$\frac{\partial N}{\partial t} + \frac{\partial (N\bar{a}')}{\partial a} + \frac{\partial (N\bar{b}')}{\partial b} + \frac{\partial (N\bar{c}')}{\partial c} = 0. \quad (1)$$

D'une manière analogue on pourra déduire de l'état du système au moment t le nombre total des points qui se trouvent, au bout du temps $t + \tau$, à l'intérieur du parallélipède $d\omega$, si en même temps les composantes de vitesse sont comprises dans les limites du parallélipède $d\omega'$; c'est-à-dire qu'on doit avoir

$$N(t + \tau, a, b, c) s(t + \tau, a, b, c, a', b', c') \\ = \iiint dx'' dy'' dz'' N(t, x, y, z) s s_1,$$

où s dépend des variables t, x, y, z, x', y', z' et s_1 de ces mêmes quantités et, en outre, de $x'' y'' z''$, variables qui satisferont aux conditions suivantes

$$a = x + x'\tau, \quad b = y + y'\tau, \quad c = z + z'\tau, \\ a' = x' + x''\tau, \quad b' = y' + y''\tau, \quad c' = z' + z''\tau,$$

les puissances supérieures de τ étant négligées. On remplacera donc partout ci-dessus x par $a - a'\tau$, x' par $a' - x''\tau$ et ainsi de suite, et les intégrations seront étendues à toutes les valeurs possibles de x'', y'', z'' . On en déduira, par comparaison des coefficients de τ ,

$$\left. \begin{aligned} & \frac{\partial (Ns)}{\partial t} + \frac{\partial (Ns)}{\partial a} a' + \frac{\partial (Ns)}{\partial b} b' + \frac{\partial (Ns)}{\partial c} c' \\ & + N \left(\frac{\partial (\overline{sa''})}{\partial a'} + \frac{\partial (\overline{sb''})}{\partial b'} + \frac{\partial (\overline{sc''})}{\partial c'} \right) = 0, \end{aligned} \right\} \quad (2)$$

en écrivant N au lieu de $N(t, a, b, c)$ et s au lieu de $s(t, a, b, c, a', b', c')$, tandis que $\bar{a}'', \bar{b}'', \bar{c}''$ sont les valeurs moyennes de x'', y'', z'' qui correspondent au point a, b, c et aux composantes de vitesse a', b', c' , c'est à dire

$$\overline{a''} = \iiint dx'' dy'' dz'' s_1(t, a, b, c, a', b', c', x'' y'' z'') x''$$

et ainsi de suite.

J'appellerai les équations (1) et (2) *première* et *seconde équation de continuité*. Ce sont les premières de toute une suite d'équations de continuité qui peuvent facilement être formées par le procédé employé ici. Chaque équation de la suite peut sans difficulté être déduite de la suivante. C'est ainsi qu'on retrouvera la première équation de continuité en multipliant la seconde par $da'db'dc'$ et intégrant entre les limites des variables. Pour ces limites s s'évanouira, et l'on aura de plus

$$\iiint da' db' dc' s(a, b, c, a', b', c') = 1,$$

si les variables parcourent toutes les valeurs possibles. En général, on peut de cette manière déduire les première, seconde,, $(n-1)^{\text{ième}}$ équations de la $n^{\text{ième}}$ qui, par suite, peut être considérée comme une généralisation de toutes les équations précédentes.

Ces équations purement cinétiques, peuvent être employées dans la dynamique de la manière suivante. Imaginons que tous les points du système aient la même masse m , qui est supposée être une quantité invariable comme le nombre de ces points matériels, de manière qu'aucun d'eux ne puisse être divisé en plusieurs et que plusieurs ne puissent être réunis en un.

En général, tout corps est supposé formé de plusieurs systèmes de cette espèce. On ne fait de la sorte aucune hypothèse sur la nature du corps, sauf celle de l'invariabilité de sa masse; supposer le nombre des points invariable au cours des calculs, n'est qu'une fiction purement mathématique sans conséquence physique. Les

équations de continuité établies ci-dessus sont valables pour chaque système en particulier.

De plus toute pression exercée sur un élément superficiel de l'espace peut être considérée comme produite par le mouvement des points matériels qui passent par cet élément, la pression suivant une direction donnée étant définie avec plus de précision comme la quantité de mouvement, estimée suivant la direction donnée, des points qui dans l'unité de temps passent par l'élément considéré. Dans le cas d'une pression variable, cette quantité de mouvement doit être mesurée par la quantité de mouvement au moment considéré. Cette définition est une conséquence directe de la décomposition supposée du corps en points matériels mobiles, et elle aussi n'introduit aucune hypothèse physique.

La pression sur une surface en mouvement peut être conçue d'une manière analogue: il suffit d'imaginer que la surface se meuve en même temps que le corps, de sorte que tout élément de la surface par lequel passe un grand nombre de points matériels reçoive le mouvement moyen de ces points. Par „mouvement moyen“ on entend un mouvement dont les vitesses composantes suivant les trois axes sont les vitesses des éléments perpendiculaires à ces trois axes se mouvant de manière à faire passer des masses égales des deux côtés.

Si l'on désigne par Σ une somme étendue à tous les systèmes de points matériels, dont est composé le corps, on verra, en employant la notation ci-dessus, que

$$\Sigma m N \overline{a'} dt db dc$$

est la masse totale qui dans l'élément de temps dt passe par la surface latérale $db dc$ du parallélépipède. Si cette

surface a une vitesse telle qu'aucune masse n'y passe, on aura

$$\Sigma m N(\bar{a}' - u) = 0.$$

u est donc la composante de la vitesse du mouvement moyen dans la direction de l'axe des x . Les deux autres composantes sont déterminées par

$$\Sigma m N(\bar{b}' - v) = 0, \quad \Sigma m N(\bar{c}' - w) = 0.$$

Si l'on multiplie l'équation (1) par m , et si l'on forme les équations correspondantes pour tous les systèmes, puis qu'on additionne toutes ces équations, on trouvera

$$\Sigma \left[\frac{\partial(mN)}{\partial t} + \frac{\partial(mN\bar{a}')}{\partial a} + \frac{\partial(mN\bar{b}')}{\partial b} + \frac{\partial(mN\bar{c}')}{\partial c} \right] = 0.$$

Cette équation, combinée avec les équations ci-dessus, donnera

$$\frac{\partial \rho}{\partial t} + \frac{\partial(\rho u)}{\partial a} + \frac{\partial(\rho v)}{\partial b} + \frac{\partial(\rho w)}{\partial c} = 0, \quad (3)$$

où l'on a remplacé $\Sigma m N$ par ρ , qui représente en conséquence la somme des masses des points matériels compris dans l'unité de volume; par suite ρ n'est autre chose que la densité du corps au point a, b, c . De plus, comme u, v, w sont les composantes de la vitesse de ce même point du corps, si l'on ne tient compte, comme à l'ordinaire, que du mouvement moyen du corps, l'équation (3) est tout à fait identique à l'équation bien connue qu'on nomme spécialement *l'équation de continuité*.

Si l'on multiplie l'équation (2) par $ma'da'db'dc'$ et qu'on l'intègre entre les limites des variables, puis qu'on additionne les équations correspondantes pour tous les systèmes, on obtiendra

$$\Sigma \left[\frac{\partial(mN\bar{a}')}{\partial t} + \frac{\partial(mN\bar{a}'^2)}{\partial a} + \frac{\partial(mN\bar{a}'b')}{\partial b} + \frac{\partial(mN\bar{a}'c')}{\partial c} - mN\bar{a}'' \right] = 0. \quad (4)$$

Ici $\overline{a'}$ désigne la valeur moyenne de $\overline{a''}$ pour toutes les valeurs de a', b', c' . En ce qui concerne le calcul indiqué, on doit remarquer que l'on a effectué l'intégration par parties des trois derniers termes de l'équation (2) et que les intégrales qui correspondent aux limites de a', b', c' s'évanouissent: en effet s entre comme facteur dans ces intégrales, et les limites de a', b', c' sont choisies de manière que la probabilité s s'évanouisse pour ces limites, qui peuvent être considérées comme finies, même si elles sont aussi grandes qu'on veut.

Par suite si P_1, T_3, T_2 sont les pressions dans les directions des axes des x , des y et des z sur l'unité de surface du corps sur la face latérale $db\,dc$ du parallélépipède $d\omega$, on aura, en conséquence de la définition donnée,

$$P_1 = \Sigma m N(\overline{a' - u})^2, \quad T_3 = \Sigma m N(\overline{a' - u})(\overline{b' - v}), \\ T_2 = \Sigma m N(\overline{a' - u})(\overline{c' - w}),$$

d'où, en ayant égard aux équations

$$\Sigma m N(\overline{a' - u}) = 0, \quad \Sigma m N(\overline{b' - v}) = 0, \quad \Sigma m N(\overline{c' - w}) = 0,$$

on déduira

$$\Sigma m N \overline{a'^2} = P_1 + \rho u^2, \quad \Sigma m N \overline{a'b'} = T_3 + \rho uv, \\ \Sigma m N \overline{a'c'} = T_2 + \rho uw.$$

Si l'on porte ces valeurs dans l'équation (4), et qu'on en retranche l'équation (3) multipliée par u , on obtiendra

$$\rho \left(\frac{\partial u}{\partial t} + \frac{\partial u}{\partial a} u + \frac{\partial u}{\partial b} v + \frac{\partial u}{\partial c} w \right) + \frac{\partial P_1}{\partial a} + \frac{\partial T_3}{\partial b} + \frac{\partial T_2}{\partial c} = \rho A, \quad (5)$$

où ρA remplace $\Sigma m N \overline{a''}$. A peut être considéré comme l'accélération du mouvement moyen au point a, b, c dans la direction de l'axe des x .

Des expressions analogues peuvent être formées pour les autres faces latérales du parallépipède $d\omega$.

Ces équations sont bien connues et sont employées pour le calcul des mouvements et des pressions dans les corps, tant solides et élastiques que liquides; on y introduit seulement, selon la nature du problème, certaines suppositions restrictives. Au contraire, les équations qui dépendent seulement des mouvements moyens ne suffisent pas à résoudre tous les problèmes proposés par la science moderne, comme la détermination de la diffusion, le frottement intérieur, la propagation de la force vive, qui est aussi intimement liée à la conduction de la chaleur.

Ce sont de tels problèmes qui indiquent la marche suivie ici. Les équations de condition établies ici sont les équations cinétiques fondamentales; car elles expriment les lois générales du mouvement dans l'intérieur des corps, indépendamment des forces mises en jeu, *et elles sont en même temps*, si l'on passe de la manière indiquée aux pressions intérieures des corps, *une généralisation des équations dynamiques ordinaires* (équation (5) et les deux analogues). C'est pourquoi elles peuvent être considérées comme servant d'introduction à la discussion générale des problèmes mentionnés plus haut, tandis que la solution de ces problèmes ne peut, cela va sans dire, être obtenue que par l'introduction d'hypothèses particulières.

SUR LE DÉVELOPPEMENT
DES FONCTIONS ARBITRAIRES
AU MOYEN DE FONCTIONS DONNÉES.

SUR LE DÉVELOPPEMENT DES FONCTIONS ARBITRAIRES AU MOYEN DE FONCTIONS DONNÉES.

TIDSSKRIFT FOR MATEMATIK 1876, P. 129—144.

Le développement d'une fonction donnée $f(x)$ en série infinie ou par une intégrale définie où la variable n'entre que dans des fonctions données peut être rapporté à la forme

$$f(x) = \sum_k A_k F(x, k). \quad (1)$$

Le procédé pour la détermination d'un tel développement consiste généralement à déterminer d'abord les coefficients A_k du développement, supposé possible, puis à discuter *les conditions de validité de ce développement*. Nous supposons que les coefficients sont déterminés par

$$A_k = \int_a^b f(x') \varphi(x', k) dx', \quad (2)$$

et que l'intégration est effectuée suivant un chemin réel, de a à b . On reconnaît alors immédiatement que le développement trouvé n'est valable que si x est réel et compris entre les limites a et b ; car le second membre de l'expression (1) est indépendant de toutes les valeurs que peut prendre la fonction $f(x)$, quand la variable est imaginaire ou quand elle reste en dehors desdites limites. Après avoir déterminé la fonction $\varphi(x)$, on doit discuter plus en détail les conditions de validité du développe-

ment et cela se fait généralement d'après une méthode que Dirichlet a imaginée à propos des développements au moyen des fonctions circulaires (sin. et cos.) et au moyen des fonctions sphériques (Repert. de Dove, tome 1 et Journal de Crelle, tome 17). La marche de cette méthode est de chercher d'abord la somme d'un nombre fini de termes de la série (1), puis d'exécuter l'intégration indiquée dans (2) et finalement de chercher la valeur limite de la somme, quand le nombre des termes croît à l'infini. Si l'on cherche cette valeur limite avant d'effectuer les intégrations, le résultat sera en général indéterminé, bien que fini; aussi les formules exigent elles que les intégrations soient exécutées d'abord, et, en général, il n'est pas permis d'intervertir l'ordre; pourtant, il va sans dire qu'on a le droit, comme le fait Dirichlet, de sommer d'abord un nombre arbitraire, mais fini, de termes, avant d'effectuer l'intégration.

C'est cette marche qu'on a suivie de préférence, après lui. Ainsi on la retrouvera dans le „Handbuch der Kugelfunctionen“ de Heine, dans les „Partielle Differentialgleichungen“ de Riemann, dans le mémoire de Hankel „Die Fourier'schen Reihen“. Dans ce dernier mémoire (Mathem. Ann., tome 8) la méthode est, en tant que le permet la nature du problème, appliquée pour la première fois au développement par les fonctions cylindriques ou fonctions besséliennes (de la première espèce). Dans un mémoire de Schläfli, récemment publié (Mathem. Ann., tome 10) le même problème est traité d'une manière nouvelle; mais la méthode est encore celle de Dirichlet.

Sans doute la théorie du développement des fonctions arbitraires au moyen de fonctions données est

d'une grande importance, surtout pour la physique mathématique; aussi est-il à prévoir qu'elle s'étendra successivement à des fonctions de développement beaucoup plus nombreuses que celles qu'on a employées jusqu'ici. Mais la méthode de Dirichlet comporte déjà dans les cas traités par lui-même des difficultés assez considérables, qui proviennent en premier lieu de ce qu'on doit sommer un nombre fini de termes d'une série non convergente et puis intégrer entre des limites finies le résultat trouvé. Mais cette difficulté n'est pas inhérente à la nature du problème. Le problème exige, comme on voit, d'abord une intégration par rapport à x dans (2), et cette opération peut toujours être facilement exécutée, au moins entre deux limites infiniment rapprochées. Puis on peut effectuer la sommation (1), pourvu qu'on se borne à cet élément de l'intégrale; car la série doit toujours à présent être convergente, et en traitant de la même manière tous les éléments dans lesquels l'intégrale (2) peut être décomposée, on obtiendra la somme cherchée.

Pour opérer ladite intégration entre les limites x_1 et x_2 , on pourra appliquer la proposition suivante

$$\left. \begin{aligned} & \int_{x_1}^{x_2} dx' \int_{x_1}^{x_2} dx'' (f(x') - f(x'')) (\varphi(x') - \varphi(x'')) \\ &= 2(x_2 - x_1) \int_{x_1}^{x_2} f(x') \varphi(x') dx' - 2 \int_{x_1}^{x_2} f^2(x') dx' \int_{x_1}^{x_2} \varphi(x') dx'. \end{aligned} \right\} (3)$$

Ici l'on suppose que $x_2 - x_1$ est une quantité infiniment petite. Si les deux fonctions f et φ sont finies et continues entre les limites x_1 et x_2 , on reconnaît facilement que le premier membre de l'équation est infiniment petit de l'ordre de $(x_2 - x_1)^4$, tandis que les deux termes

du second membre sont chacun de l'ordre de $(x_2 - x_1)^2$. Dans ce cas, où l'on peut remplacer $\frac{1}{x_2 - x_1} \int_{x_1}^{x_2} f(x') dx'$ par $f(x_1)$ ou $f(x_2)$ on aura donc

$$\int_{x_1}^{x_2} f(x') \varphi(x') dx' = f(x_1) \int_{x_1}^{x_2} \varphi(x') dx'. \quad (4)$$

Dans le cas où l'une des deux fonctions, $\varphi(x)$ par exemple, est finie mais varie infiniment vite, de manière qu'elle est comparable à $\cos kx$ ou $\sin kx$ pour k infiniment grand, $\varphi(x') - \varphi(x'')$ aura une valeur finie et le premier membre de (3) sera de l'ordre de $\frac{(x_2 - x_1)^2}{k}$, tandis que chacun des deux termes du second membre

* NOTE 1. sera de l'ordre de $\frac{x_2 - x_1}{k}$.* Donc l'équation (4) sera encore valable dans ce cas, qui se présentera précisément dans ce qui suit. Au contraire, si les deux fonctions $\varphi(x)$ et $f(x)$ sont de cette espèce, tous les termes de l'équation (3) seront du même ordre et l'équation (4) ne sera plus valable.

Donc si l'intégrale (2) est décomposée en une somme d'intégrales de la même espèce que le premier membre de l'équation (4) où les deux limites de chaque intégrale sont infiniment rapprochées, l'équation (4) sera valable pour tous les éléments d'intégrale, pourvu que $f(x)$ soit une fonction finie et continue entre les limites a et b et que $\varphi(x)$ soit finie et continue ou varie infiniment vite. Mais, $f(x)$ étant finie et discontinue, si l'on a soin que les points x_1 et x_2 coïncident avec les points de discontinuité de la fonction, ce qui est toujours possible, à condition que la discontinuité ne se produise qu'en un

nombre fini de points entre a et b , l'équation (4) sera encore applicable dans ce cas.

Si la fonction $f(x)$ devient infinie en un seul point ou en un nombre fini de points c , x_1 ou x_2 , devra de même coïncider avec l'un de ces points; par suite, l'un des éléments de l'intégrale desquels on tient compte dans la somme (1) sera

$$\sum_k \int_{x_1}^c f(x') \varphi(x', k) F(x, k) dx'. \quad (5)$$

Ici x_1 peut être suffisamment voisin de c pour que $f(x')$ soit toujours ou croissante ou décroissante entre les limites de l'intégrale. Alors, d'après une proposition employée par Dirichlet, l'expression ci-dessus sera comprise entre les limites

$$M \int_{x_1}^c f(x') dx' \quad \text{et} \quad N \int_{x_1}^c f(x') dx,$$

M et N étant la plus grande et la plus petite des valeurs que peut prendre

$$\sum_k \varphi(x', k) F(x, k)$$

entre les limites de l'intégrale. Donc, si cette quantité est finie, même si elle est indéterminée, et si $f(x)$ devient infinie de telle manière que $\int_a^x f(x') dx'$ reste finie et continue pour toute valeur de x comprise entre a et b , la somme (5) convergera vers zéro.

Après avoir déterminé les éléments d'intégrale au moyen de l'équation (4), on doit opérer la sommation par rapport à k indiquée dans l'équation (1). Dans une classe de développements fonctionnels qui sont d'une grande importance pour la physique mathématique, les

k sont les racines de l'équation $F(b, k) = 0$, et le développement entraîne que $f(x)$ soit nulle pour la valeur limite $x = b$.

C'est pourquoi j'exposerai quelques formules générales relatives à la sommation de cette espèce de séries.

Soit donnée une fraction $\frac{p(y)}{q(y)}$, dont le numérateur et le dénominateur peuvent tous deux être développés en série suivant les puissances entières et positives de y , ces séries étant convergentes pour toutes les valeurs de y . De plus, soient k les racines de l'équation $q(y) = 0$; on suppose qu'aucune de ces racines n'annule $p(y)$. On considère l'intégrale

$$\frac{1}{2\pi i} \int \frac{p(t)}{q(t)} \cdot \frac{dt}{t-y} = Q, \quad (6)$$

où la variable t est supposée complexe. Cette variable est supposée parcourir dans le sens positif (sens inverse de celui des aiguilles d'une montre) un cercle de rayon ρ et ayant son centre au point origine. Donc si l'on pose $t = \rho e^{\theta i}$, on aura

$$Q = \frac{1}{2\pi} \int_0^{2\pi} \frac{p(\rho e^{\theta i})}{q(\rho e^{\theta i})} \cdot \frac{\rho e^{\theta i} d\theta}{\rho e^{\theta i} - y}. \quad (7)$$

ρ croissant, Q convergera vers zéro, si l'on a

$$\left(\frac{p(\rho e^{\theta i})}{q(\rho e^{\theta i})} \right)^{\rho \rightarrow \infty} = \begin{cases} 0, & \text{pour } \sin \theta \geq 0, \\ a, & \text{pour } \sin \theta = 0, \end{cases} \quad (8)$$

où a désigne une quantité finie.

D'autre part, on peut déterminer l'intégrale Q en décrivant dans le sens positif des cercles infiniment petits autour des points qui correspondent aux racines de

l'équation $q(t)(t-y) = 0$. De celle manière, si toutes les racines k sont différentes les unes des autres et différentes de y , on remplacera l'équation (6) par celle-ci

$$Q = \frac{p(y)}{q(y)} + \sum_k \frac{p(k)}{q'(k)(k-y)}. \quad (9)$$

Par suite, on peut égaler cette expression à (7), pourvu que la sommation soit étendue à toutes les racines k plus petites que ρ et que de même y soit plus petit que ρ^* . ρ croissant à l'infini, on aura donc, * NOTE 2 si la condition (8) est remplie,

$$\frac{p(y)}{q(y)} + \sum_k \frac{p(k)}{q'(k)(k-y)} = 0, \quad (10)$$

la sommation étant étendue à toutes les racines k (qui sont en nombre infini) rangées d'après l'ordre de grandeur de leurs modules. On voit ainsi que la règle de décomposition des fractions rationnelles peut encore être appliquée à l'espèce de fractions considérée ici, pourvu que la condition (8) soit remplie.

Si le développement de $q(y)$ suivant les puissances croissantes de y commence par une puissance positive de y , comme y^m , et si celui de $p(y)$ commence par la même puissance ou par une puissance supérieure, on pourra, pour éviter les racines $k = 0$ communes au numérateur et au dénominateur, diviser p et q par y^m ; mais on reconnaît facilement que la formule (10) sera encore valable, si l'on ne fait pas cette réduction, à condition qu'on ne tienne pas compte des racines $k = 0$ dans la sommation.

On peut encore faire ce calcul dans le cas où l'équation $q(y) = 0$ a des racines égales.

Nous considérerons en particulier le cas de $\frac{p(y)}{(q(y))^2}$, où, comme ci-dessus, $q(y) = 0$ n'admet pas de racines égales ni de racines qui satisfassent à $p(y) = 0$. Alors, par une analyse analogue à la précédente, ou en appliquant les règles de la décomposition des fractions, on trouvera

$$\frac{p(y)}{(q(y))^2} + \sum_k \frac{1}{q'(k)} \frac{d}{dk} \left(\frac{p(k)}{q'(k)(k-y)} \right) = 0, \quad (11)$$

pourvu que la conditions correspondante à (8)

$$\left(\frac{p(\rho e^{i\theta})}{(q(\rho e^{i\theta}))^2} \right)^{\rho=\infty} = \begin{cases} 0, & \sin \theta \geq 0, \\ a, & \sin \theta = 0, \end{cases} \quad (12)$$

soit remplie. Si le développement de $q(y)$ commence par y^m et celui de $p(y)$ par y^{2m} ou par une puissance supérieure, on peut employer la même formule, à condition qu'on ait soin de négliger les racines $k = 0$.

Je vais maintenant faire l'application de la méthode générale indiquée ici à deux développements d'une fonction arbitraire au moyen des fonctions besséliennes de la première espèce. Ces fonctions sont définies pour tout indice n , entier et positif, par l'égalité

$$J_n(r) = \left\{ \begin{aligned} & \left(\frac{r}{2} \right)^n \cdot \frac{1}{n!} - \left(\frac{r}{2} \right)^{n+2} \cdot \frac{1}{(n+1)! 1!} \\ & + \left(\frac{r}{2} \right)^{n+4} \cdot \frac{1}{(n+2)! 2!} - \dots, \end{aligned} \right\} \quad (13)$$

d'où l'on déduit, pour $n = 0$,

$$J_0(r) = 1 - \frac{r^2}{2^2} + \frac{r^4}{2^2 4^2} - \frac{r^6}{2^2 4^2 6^2} + \dots$$

Nous appellerons l'attention sur les équations sui-

vantes que nous emploierons ultérieurement et qui peuvent facilement être déduites de (13)

$$\frac{d(r^{-n}J_n(r))}{dr} = -r^{-n}J_{n+1}(r), \quad (14)$$

$$\frac{d(r^n J_n(r))}{dr} = r^n J_{n-1}(r), \quad (15)$$

$$\frac{1}{r} \frac{d}{dr} \left(r \frac{dJ_n(r)}{dr} \right) + \left(1 - \frac{n^2}{r^2} \right) J_n(r) = 0. \quad (16)$$

Si l'on admet que les équations (14) et (15) sont encore valables pour $(n-1)$ négatif, on déduira successivement de ces équations

$$\begin{aligned} J_{-1}(r) &= -J_1(r), \quad J_{-2}(r) = J_2(r), \\ J_{-n}(r) &= (-1)^n J_n(r). \end{aligned}$$

L'extension de la définition au cas des indices fractionnaires ne devant être d'aucune utilité dans ce qui suivra, nous l'omettons. Nous ajouterons encore que $J_n(r)$ est une fonction périodique qui, pour des valeurs croissantes de r , s'approche de $\sqrt{\frac{2}{\pi r}} \cos\left(r - \frac{2n+1}{4} \pi\right)$, et que l'équation $J_n(r) = 0$ n'a que des racines réelles: pour la démonstration de ces propriétés, nous renverrons aux „Studien über die Bessel'schen Funktionen“ de Lommel.*

* NOTE 3.

En conséquence des équations (16), nous aurons

$$\begin{aligned} \frac{1}{r} \frac{d}{dr} \left(r \frac{dJ_n(ar)}{dr} \right) + \left(a^2 - \frac{n^2}{r^2} \right) J_n(ar) &= 0, \\ \frac{1}{r} \frac{d}{dr} \left(r \frac{dJ_n(\beta r)}{dr} \right) + \left(\beta^2 - \frac{n^2}{r^2} \right) J_n(\beta r) &= 0. \end{aligned}$$

Si l'on multiplie la première de ces relations, par $J_n(\beta r)$ et la seconde par $J_n(ar)$, on obtiendra par soustraction

$$\frac{1}{r} \frac{d}{dr} \left(r \left(J_n(\beta r) \frac{dJ_n(\alpha r)}{dr} - J_n(\alpha r) \frac{dJ_n(\beta r)}{dr} \right) \right) + (\alpha^2 - \beta^2) J_n(\alpha r) J_n(\beta r) = 0.$$

Cette équation, multipliée par $r dr$ et intégrée entre les limites $r = 0$ et $r = 1$, donnera

$$\left. \begin{aligned} & \beta J_n(\alpha) J'_n(\beta) - \alpha J_n(\beta) J'_n(\alpha) \\ &= (\alpha^2 - \beta^2) \int_0^1 J_n(\alpha r) J_n(\beta r) r dr, \end{aligned} \right\} \quad (17)$$

où J'_n représente, suivant la notation de Lagrange, la dérivée de J_n .

Si nous supposons que k et k_1 sont des racines différentes de l'équation $J_n(r) = 0$, on aura

$$\int_0^1 J_n(kr) J_n(k_1 r) r dr = 0, \quad (18)$$

et si k' et k'_1 sont de racines différentes de la dérivée $J'_n(r)$, l'équation (17) donnera de même

$$\int_0^1 J_n(k' r) J_n(k'_1 r) r dr = 0. \quad (19)$$

Différentions l'équation (17) par rapport à β et faisons ensuite $\alpha = \beta$; en ayant égard à l'équation (16), on trouvera

$$2 \int_0^1 (J_n(\alpha r))^2 r dr = \left(1 - \frac{n^2}{\alpha^2} \right) (J_n(\alpha))^2 + (J'_n(\alpha))^2,$$

et par suite

$$2 \int_0^1 (J_n(kr))^2 r dr = (J'_n(k))^2, \quad (20)$$

$$2 \int_0^1 (J_n(k' r))^2 r dr = \left(1 - \frac{n^2}{k'^2} \right) (J_n(k'))^2. \quad (21)$$

Nous considérerons à présent le développement traité par Hankel et Schäfli

$$f(r) = \sum_k A_k J_n(kr), \quad (22)$$

où la sommation est étendue à toutes les racines positives k_1 de $J_n(r) = 0$. D'après ce qui précède, on pourra déterminer les coefficients A_k en multipliant l'équation par $2J_n(k_1 r) r dr$ et intégrant de $r = 0$ à $r = 1$; par où l'on obtiendra, en vertu des équations (18) et (20)

$$2 \int_0^1 J_n(k_1 r) r f(r) dr = A_k (J'_n(k_1))^2,$$

ou, en écrivant r' au lieu de r et k au lieu de k_1 ,

$$A_k = \frac{2}{(J'_n(k))^2} \int_0^1 J_n(kr') r' f(r') dr'. \quad (23)$$

Il suit de là que le développement (22) déterminé de cette manière ne peut être valable que pour r réel et compris entre 0 et 1. Pour ces limites elles-mêmes, le second membre de (22) s'évanouira, sauf dans le cas $n = 0$ et $r = 0$. Pour discuter avec plus de précision les conditions de validité du développement, il faut considérer la somme

$$\sum_k \frac{2 J_n(kr)}{(J'_n(k))^2} \int_{r_1}^{r_2} J_n(kr') r' f(r') dr', \quad (24)$$

où les limites r_1 et r_2 de l'intégrale sont comprises entre 0 et 1 et où $r_2 - r_1$ est une quantité infiniment petite et positive.

Pour exécuter la sommation qui se présente ici, on fera d'abord dans (11) $q(y) = J_n(y)$ et $y = 0$; comme $\frac{d}{dk} (k J'_n(k)) = 0$, on obtiendra ainsi, d'après (16),

$$\left[\frac{p(y)}{(J_n(y))^2} \right]^{y=0} = - \sum_k \frac{p'(k)}{(J_n'(k))^2 k}. \quad (25)$$

Si l'on fait de plus $p(y) = J_{n+m}(yr) J_{n+m}(yr_1)$, où m est un entier positif ou 0, on trouvera sans difficulté que la condition (12) est remplie si la somme de r et r_1 , qui

* NOTE 4. tous deux sont positifs, est plus petite que 2^* , et comme on obtiendra au moyen de (14) et (15)

$$p'(k) = r J_{n+m-1}(kr) J_{n+m}(kr_1) - r_1 J_{n+m}(kr) J_{n+m+1}(kr_1),$$

l'équation (25) donnera pour $m > 0$

$$\sum \frac{J_{n+m-1}(kr) J_{n+m}(kr_1)}{(J_n'(k))^2 k} = \frac{r_1}{r} \sum_k \frac{J_{n+m}(kr) J_{n+m+1}(kr_1)}{(J_n'(k))^2 k}.$$

Ces séries sont convergentes et les sommes conservent pour des m croissants une valeur finie, ce qu'on peut reconnaître en introduisant les valeurs limites des fonc-

* NOTE 5. tions besséliennes pour k croissant à l'infini*. Si donc on fait successivement $m = 1, 2, 3, \dots$, et si l'on multiplie les équations obtenues de cette manière, on reconnaît que pour $r_1 < r$

$$\sum_k \frac{J_n(kr) J_{n+1}(kr_1)}{(J_n'(k))^2 k} = 0, \quad (r_1 < r). \quad (26)$$

Pour $m = 0$, le premier membre de (25) deviendra $r^n r_1^n$ et l'on aura, à cause du résultat ci-dessus,

$$2 \sum_k \frac{J_{n-1}(kr) J_n(kr_1)}{(J_n'(k))^2 k} = - r^{n-1} r_1^n, \quad r_1 < r. \quad (27)$$

La somme est doublée ici, parce que le même terme se répète pour deux racines égales, mais de signes contraires; elle ne s'étend donc qu'aux racines positives.

Ensuite, pour déterminer les mêmes sommes dans

le cas de r_1 égal ou supérieur que r , nous considérerons l'expression

$$u = 2 \sum_k \frac{J_n(kr) J_n(kr_1)}{(J'_n(k))^2 k^2}.$$

En employant les formules (14) et (15) on en déduira

$$\begin{aligned} \frac{d(ur_1^{-n})}{dr_1} &= -2r_1^{-n} \sum_k \frac{J_n(kr) J_{n+1}(kr_1)}{(J'_n(k))^2 k}, \\ \frac{d(ur^n)}{dr} &= r^n \sum_k \frac{J_{n-1}(kr) J_n(kr_1)}{(J'_n(k))^2 k}. \end{aligned}$$

Si nous avons comme ci-dessus $r_1 < r$, on verra, d'après (26) et (27), que la première de ces expressions s'évanouit; la seconde devient $-r^{2n-1}r_1^n$. Il en résulte que

$$u = \frac{r_1^n}{2n} (cr^{-n} - r^n)$$

où c est une constante provisoirement indéterminée. On aura de plus

$$\frac{d(ur_1^n)}{dr_1} = 2r_1^n \sum_k \frac{J_n(kr) J_{n-1}(kr_1)}{(J'_n(k))^2 k} = r_1^{2n-1} (cr^{-n} - r^n).$$

Toutes ces équations qui sont déduites de (26) et (27) sont, comme ces relations elles-mêmes, valables pour $r_1 < r$ et $r + r_1 < 2$. On peut donc poser $r = 1$, si r_1 reste plus petit que l'unité, et si l'on porte cette valeur de r dans la dernière équation, $J_n(k)$ étant nul, on trouvera $c = 1$. Si l'on tient compte de cette valeur de c , la même équation donnera, par permutation de r avec r_1 ,

$$2 \sum_k \frac{J_{n-1}(kr) J_n(kr_1)}{(J'_n(k))^2 k} = r^{n-1}(r_1^{-n} - r_1^n), \quad (r_1 > r). \quad (28)$$

Enfin, si dans la relation

$$\frac{d(u r^{-n})}{dr} = -2 r^{-n} \sum_k \frac{J_{n+1}(kr) J_n(kr_1)}{(J'_n(k))^2 k},$$

on introduit la valeur trouvée de u et qu'on permute r et r_1 , on obtiendra

$$2 \sum_k \frac{J_n(kr) J_{n+1}(kr_1)}{(J'_n(k))^2 k} = r^n r_1^{-n-1}, \quad (r_1 > r). \quad (29)$$

La fonction u elle-même peut facilement être déterminée, si $r_1 > r$, en permutant r_1 et r . On verra alors que cette fonction qui est continue, acquiert pour $r_1 = r$ la valeur $\frac{r^n}{2n} (r^{-n} - r^n)$, sauf dans le cas limite $r = 0$. Par conséquent, on aura

$$2 \sum_k \frac{(J_n(kr))^2}{(J'_n(k))^2 k^2} = \frac{r^n}{2n} (r^{-n} - r^n),$$

équation qui, multipliée successivement par r^{2n} et par r^{-2n} , donnera, après différentiation par rapport à r ,

$$2 \sum_k \frac{J_{n-1}(kr) J_n(kr)}{(J'_n(k))^2 k} = \frac{1}{2r} r^{2n-1}, \quad (30)$$

$$2 \sum_k \frac{J_n(kr) J_{n+1}(kr)}{(J'_n(k))^2 k} = \frac{1}{2r}, \quad (31)$$

pour $0 < r < 1$.

Nous revenons maintenant à la sommation qui figure dans l'expression (24). D'après l'équation (4) et sous les conditions relatives à cette équation, on pose

$$\int_{r_1}^{r_2} J_n(kr') r' f(r') dr' = f(r_1) r_1^n \int_{r_1}^{r_2} J_n(kr') r'^{1-n} dr';$$

après l'intégration faite, cette expression deviendra

$$- \frac{f(r_1) r_1^n}{k} [r_2^{1-n} J_{n-1}(kr_2) - r_1^{1-n} J_{n-1}(kr_1)], \quad (32)$$

où l'on pourrait d'ailleurs dans l'expression $f(r_1)r_1^n$ remplacer r_1 par r_2 ou par une quantité arbitraire entre r_1 et r_2 . Il reste, pour obtenir l'expression (24), à calculer les deux sommes

$$r_2^{1-n} \sum_k \frac{2 J_n(kr) J_{n-1}(kr_2)}{(J'_n(k))^2 k} \quad \text{et} \quad r_1^{1-n} \sum_k \frac{2 J_n(kr) J_{n-1}(kr_1)}{(J'_n(k))^2 k}.$$

D'après (28), ces deux sommes sont pour $r > r_2$ égales à $r^{-n} - r_2^{-n}$; de même d'après (27) elles sont égales, si $r < r_1$, à r^{-n} . Au contraire, si r est compris entre r_2 et r_1 , leur différence sera $-r^{-n}$. On voit par là que l'expression (24) s'évanouit, à moins que r soit compris entre les limites r_1 et r_2 , auquel cas elle est égale à $f(r)$. Donc le développement (22) est valable pour toute fonction $f(r)$ finie et continue entre les limites 0 et 1.

Si $f(r)$ est une fonction discontinue pour un nombre fini de valeurs de r , on fait coïncider, comme nous l'avons dit ci-dessus, les limites r_1 ou r_2 avec les points de discontinuité, lorsque ceux-ci sont infiniment rapprochés de ces limites. Si de plus r tombe en un point de discontinuité, on verra d'après (30) que, pour $r = r_1$ ou $r = r_2$, l'expression (24) est égale à $\frac{1}{2}f(r)$, ce qui montre que le développement (22) donne, pour les points de discontinuité, la valeur moyenne des deux valeurs de $f(r)$.

Si r et r_2 sont tous deux égaux à l'unité, on ne peut plus employer les formules de sommation qui supposent $r + r_2 < 2$, mais on reconnaît immédiatement que le développement (22) s'évanouit pour $r = 1$. Pour $r = r_1 = 0$, la formule de sommation (30), qu'il faudrait employer, cesse de même d'être valable, et le développement (22), ainsi que nous l'avons vu, se réduit

à zéro, à l'exception du cas où l'on a en même temps $n = 0$.

Nous étudierons le dernier cas plus en détail. Il est évident que pour $r = 0$ tous les éléments de l'intégrale (24) s'évanouiront par la sommation, à l'exception de celui qui correspond à $r_1 = 0$. Donc si l'on suppose en même temps n égal à zéro dans (23), on pourra écrire

$$A_k = \frac{2f(0)}{(J'_0(k))^2} \int_0^1 J_0(kr') r' dr' = -\frac{2f(0)}{J'_0(k)k},$$

par où le développement (22) deviendra pour $r = 0$

$$-f(0) \sum_k \frac{2}{J'_0(k)k}.$$

Cette sommation peut facilement être exécutée au moyen de la formule (10), si l'on pose $p(y) = 1$, $q(y) = J_0(y)$ et $y = 0$ ce qui donnera

$$\sum_k \frac{2}{J'_0(k)k} = -1.$$

Il en résulte que le développement (22) est encore valable dans le cas considéré.

Reste à rechercher, si le développement est encore valable quand $f(r)$ devient infinie, tandis que $\int_0^r f(r) dr$ est finie et continue. Cela dépend, ainsi que nous l'avons prouvé, de la sommation qui figure dans (22), si elle donne un résultat fini, même s'il est indéterminé, quand la sommation est exécutée avant l'intégration, r et r' étant différents. La somme à discuter est

$$\sum_k \frac{J_n(kr) J_n(kr')}{(J'_n(k))^2} \quad \text{pour } r \geq r'. \quad (33)$$

Dans cette somme $J_n(rk)$ s'approche de la valeur $\sqrt{\frac{2}{\pi kr}} \cos\left(kr - \frac{2n+1}{4}\pi\right)$, pour des valeurs croissantes de k , tandis que la $m^{\text{ième}}$ racine s'approche de la valeur $\left(m + \frac{2n-1}{4}\right)\pi$.

Si l'on compare la série ci-dessus avec

$$\sum \frac{1}{\sqrt{r'r'}} \cos\left(k_m r - \frac{2n+1}{4}\pi\right) \cos\left(k_m r' - \frac{2n+1}{4}\pi\right), \quad (34)$$

où $k_m = \left(m + \frac{2n-1}{4}\right)\pi$ et où m parcourt toutes les valeurs entières de 1 jusqu'à ∞ , on reconnaît que les termes de ces deux séries tendent à devenir les mêmes. Néanmoins il se peut que les différences des termes correspondants des deux séries ne forment pas une série convergente; aussi doit-on s'assurer du fait en développant les termes de la première série suivant les puissances croissantes de $\frac{1}{k}$. On peut d'ailleurs déterminer directement la différence des deux séries de la manière suivante.* Si l'on pose

* NOTE 6.

$$p(y) = \frac{y}{r'^2 - r^2} (r J_n(yr') J'_n(yr) - r' J_n(yr) J'_n(yr')), \quad q(y) = J_n(y),$$

$$p_1(y) = \frac{y}{\sqrt{r'r'}} \left[-r \cos\left(yr' - \frac{2n+1}{4}\pi\right) \sin\left(yr - \frac{2n+1}{4}\pi\right) + r' \cos\left(yr - \frac{2n+1}{4}\pi\right) \sin\left(yr' - \frac{2n+1}{4}\pi\right) \right],$$

$$q_1(y) = \cos\left(y - \frac{2n+1}{4}\pi\right),$$

et si l'on forme les fractions $\frac{p(y)}{(q(y))^2}$ et $\frac{p_1(y)}{(q_1(y))^2}$, bien qu'on

ne puisse pas décomposer chacune d'elles en particulier en une somme de la manière indiquée dans (11), parce que la condition (12) n'est pas remplie, on pourra pourtant appliquer l'équation (11) à la différence des deux fractions, car ici la condition (12) est remplie, ce qui est facile à reconnaître. De cette manière, on obtiendra précisément la différence des deux séries en question (33) et (34) pour le même nombre (infini) de termes, et comme chaque fraction en particulier s'évanouit pour $y = 0$, on voit que la différence cherchée est 0. Cependant la série (34) donne une somme finie, mais indéterminée; la série (33) se comportera donc de la même manière. Il suit de là que le développement (22) est valable même si $f(r)$ devient infinie, pourvu que $\int_0^r f(r) dr$ reste finie et continue.

J'appellerai finalement l'attention sur un développement nouveau au moyen des fonctions besséliennes, savoir

$$f(r) = \sum_{k'} B_{k'} J_n(k'r), \quad (35)$$

où les k' sont les racines de l'équation $J'_n(r) = 0$. Pour déterminer les coefficients $B_{k'}$, multiplions les deux membres par $J_n(k'_1 r) r dr$ et intégrons de $r = 0$ jusqu'à $r = 1$; nous obtiendrons ainsi, en vertu des équations (19) et (21)

$$2 \int_0^1 J_n(k'_1 r) r f(r) dr = B_{k'_1} \left(1 - \frac{n^2}{k'^2_1}\right) (J_n(k'_1))^2.$$

De là résulte

$$B_{k'} = \frac{2}{\left(1 - \frac{n^2}{k'^2}\right) (J_n(k'))^2} \int_0^1 J_n(k' r') r' f(r') dr'. \quad (36)$$

On peut rechercher les conditions de validité de ce développement par le même procédé qu'on a employé pour l'équation (22).^{*} C'est seulement pour $n = 0$ que le développement n'est pas valable, car alors il faut ajouter une constante C , qui n'entre pas dans la détermination de B_k , puisqu'on doit avoir

$$\int_0^1 J_0(k'r) r dr = \frac{J_1(k')}{k'} = -\frac{J'_0(k')}{k'} = 0.$$

Pour déterminer cette constante qu'on doit ajouter au second membre de (35) pour $n = 0$, on peut multiplier cette équation par $r dr$ et intégrer de $r = 0$ jusqu'à $r = 1$. On trouvera de cette manière que le développement (35) devient pour $n = 0$,

$$f(r) = 2 \int_0^1 f(r') dr' + \Sigma \frac{2 J_0(k'r)}{(J_0(k'))^2} \int_0^1 J_0(k'r') r' f(r') dr'. \quad (35')$$

Les deux développements (22) et (35) peuvent être employés dans la physique mathématique pour calculer le potentiel (par exemple, la température ou la tension électrique) dans l'intérieur d'un cylindre droit à base circulaire, quand on donne pour les deux bases le potentiel ou sa dérivée suivant la direction de l'axe (correspondant à la quantité de chaleur ou d'électricité qui entre), tandis que le potentiel ou sa dérivée suivant la direction perpendiculaire à l'axe est donné égal à zéro sur toute la surface courbe.

Dans le premier des deux cas, le potentiel peut facilement être déterminé au moyen de (22), dans le second par (35). Donc les deux développements ont des applications correspondantes et se complètent l'un par l'autre.

NOTES.

NOTE 1. Lorenz dit que les deux membres seront de l'ordre de $\frac{(x_2-x_1)^2}{k}$ et de l'ordre de $\frac{x_2-x_1}{k}$: en réalité ils sont de l'ordre de $\frac{(x_2-x_1)^3}{k}$ et de $\frac{(x_2-x_1)^2}{k}$, ce qui pourtant n'entraîne aucune modification dans les résultats.

NOTE 2. Lorenz appelle quantité complexe *plus petite que* ρ , une quantité dont le module est plus petit que ρ .

NOTE 3. La proposition concernant la réalité des racines n'est, en général, valable que si l'indice de la fonction J est un nombre entier.

NOTE 4. Cela devient évident, si l'on introduit les valeurs limites des fonctions pour une variable indépendante complexe, dont le module croît à l'infini.

NOTE 5. On reconnaît facilement que la série formée en remplaçant les fonctions besséliennes par leurs valeurs limites est convergente; mais cela n'implique pas que la série originale soit convergente. La démonstration, pour être rigoureuse, exigerait que les différences des termes correspondants des deux séries formassent une série convergente.

NOTE 6. Pour contrôler les propositions de Lorenz, nous devons chercher les valeurs des expressions $\frac{p(y)}{(q(y))^2}$ et $\frac{p_1(y)}{(q_1(y))^2}$ pour les valeurs de y dont le module est très grand.

Pour de très grands modules de y nous aurons, en employant la même proposition que ci-dessus sur la valeur limite de $J_n(y)$,

$$\begin{aligned} \frac{p(y)}{(q(y))^2} &= \frac{y}{r'^2 - r^2} \left[\frac{V \frac{\sqrt{r'}}{r} \cos\left(yr - \frac{2n+1}{4} \pi\right) \sin\left(yr' - \frac{2n+1}{4} \pi\right)}{\cos^2\left(y - \frac{2n+1}{4} \pi\right)} \right. \\ &\quad \left. - \frac{V \frac{\sqrt{r}}{r'} \cos\left(yr' - \frac{2n+1}{4} \pi\right) \sin\left(yr - \frac{2n+1}{4} \pi\right)}{\cos^2\left(y - \frac{2n+2}{4} \pi\right)} \right] \\ &= \frac{p_1(y)}{(q_1(y))^2} \frac{1}{r'^2 - r^2}. \end{aligned}$$

Si l'on pose $y = \rho e^{i\theta}$, ρ étant infiniment grand, on peut choisir ρ de manière que $\cos\left(y - \frac{2n+1}{4} \pi\right)$ ne soit égal à zéro pour aucune valeur de θ . Dans cette supposition, $\frac{p(y)}{(q(y))^2} - \frac{p_1(y)}{(q_1(y))^2} \frac{1}{r'^2 - r^2}$ sera fini pour $\sin \theta = 0$ et $\frac{p(y)}{(q(y))^2}$ s'évanouira pour toute autre valeur de θ , car l'expression $\frac{p(y)}{(q(y))^2}$ sera du même ordre de grandeur que

$$\rho e^{\rho(r+r'-2)|\sin \theta|},$$

$|\sin \theta|$ étant la valeur numérique de $\sin \theta$, et comme cette expression s'évanouit pour $\rho = +\infty$, l'expression en question s'évanouira de même.

Du reste, je suppose que Lorenz a écrit par inadvertance le dénominateur $r'^2 - r^2$ de $p(y)$ et que ce dénominateur doit être négligé.

NOTE 7. Pour examiner avec plus de précision les valeurs des séries

$$u = \sum_{k'} \frac{2 J_n(k' r) \int_0^{r_1} J_n(k' r') r' f(r') dr'}{\left(1 - \frac{n^2}{k'^2}\right) (J_n(k'))^2}$$

on a besoin de valeurs des séries

$$\sum_{k'} \frac{2 J_n(k' r_1) J_n(k' r_2)}{(k'^2 - n^2) (J_n(k'))^2}.$$

Pour trouver la somme d'une telle série, on considère l'expression $\left[\frac{p(y)(y^2 - n^2)}{y^2 (J'_n(y))^2} \right]$ en posant dans (11) $q(y) = y J'_n(y)$ et en remplaçant $p(y)$ par $p(y)(y^2 - n^2)$, ce qui donne

$$\left[\frac{p(y)(y^2 - n^2)}{y^2 (J'_n(y))^2} \right] = - \sum \frac{k'}{(k'^2 - n^2) J_n(k')} \frac{d}{dk'} \left(\frac{k' p(k')}{J_n(k') (k' - y)} \right),$$

relation où l'on doit pourtant supposer que $p(y)$ contient en facteur y^{2n} ($n > 0$) ou une puissance supérieure de y . Dans cette supposition, on aura pour $y = 0$

$$\left[\frac{p(y)(y^2 - n^2)}{y^2 (J'_n(y))^2} \right]^{y=0} = - \sum \frac{k' p'(k')}{(k'^2 - n^2) (J_n(k'))^2}. \quad (1)$$

Si l'on pose ici

$$p(y) = J_{n+m}(ry) J_{n+m}(r_1 y),$$

on aura pour $m > 0$, $n > 0$

$$- \sum \frac{k' J_{n+m-1}(k' r) J_{n+m}(k' r_1)}{(k'^2 - n^2) (J_n(k'))^2} = \frac{r_1}{r} \sum \frac{k' J_{n+m}(k' r) J_{n+m+1}(k' r_1)}{(k'^2 - n^2) (J_n(k'))^2}.$$

Ces séries sont convergentes pour $r_1 + r < 2$ et l'on verra de la même manière que plus haut (p. 506) qu'elles s'évanouissent si $r_1 < r$. Par conséquent

$$\sum_{k'} \frac{k' J_n(k'r) J_{n+1}(k'r_1)}{(k'^2 - n^2)(J_n(k'))^2} = 0 \quad (r_1 < r).$$

Si $n = 0$, $n > 0$, le premier membre de l'équation (1) devient $-r^n r_1^n$, d'où résulte

$$2 \sum_{k'} \frac{k'}{(k'^2 - n^2)} \frac{J_{n-1}(k'r) J_n(k'r_1)}{(J_n(k'))^2} = r^{n-1} r_1^n \quad (r_1 < r).$$

Par la méthode qu'emploie Lorenz on obtiendra la valeur de

$$u = 2 \sum_{k'} \frac{J_n(k'r) J_n(k'r_1)}{(k'^2 - n^2)(J_n(k'))^2}.$$

On trouve pour $r_1 < r$

$$u = \frac{r_1^n}{2n} (r^n + r^{-n}),$$

et pour $r_1 > r$

$$\begin{aligned} 2 \sum_{k'} \frac{J_n(k'r) J_{n+1}(k'r_1) k'}{(k'^2 - n^2)(J_n(k'))^2} &= r^n r_1^{n-1}, \\ 2 \sum_{k'} \frac{J_{n-1}(k'r) J_n(k'r_1) k'}{(k'^2 - n^2)(J_n(k'))^2} &= r^{n-1} (r_1^n + r_1^{-n}). \end{aligned}$$

Si l'on fait tendre r_1 vers la valeur r , on voit que u se rapprochera de $\frac{r^n}{2n} (r^n + r^{-n})$ que r_1 soit plus grand que r , ou qu'il soit plus petit. On aura donc, si $r < 1$,

$$2 \sum_{k'} \frac{(J_n(k'r))^2}{(k'^2 - n^2)(J_n(k'))^2} = \frac{r^n}{2n} (r^{-n} + r^n),$$

d'où l'on déduira facilement

$$\begin{aligned} 2 \sum_{k'} \frac{J_{n-1}(k'r) J_n(k'r) k'}{(k'^2 - n^2)(J_n(k'))^2} &= \frac{1}{2r} + r^{2n-1}, \\ 2 \sum_{k'} \frac{J_n(k'r) J_{n+1}(k'r) k'}{(n^2 - k'^2)(J_n(k'))^2} &= \frac{1}{2r}. \end{aligned}$$

A présent, on peut, de la manière développée par Lorenz, chercher la valeur de la série en question,

$$\sum_{k'} \frac{2 J_n(k' r)}{\left(1 - \frac{n^2}{k'^2}\right) J_n(k')^2} \int_0^1 J_n(k' r') r' f(r') dr'.$$

Les développements faits dans cette note ne sont pour-
tant pas valables pour $n = 0$. Il faut dans ce cas dé-
composer l'expression $\frac{p(y)}{(J'_0(y))^2}$ en fractions partielles et
chercher la valeur de l'expression pour $y = 0$. Si l'on
pose

$$p(y) = J_m(r y) J_m(r_1 y),$$

on aura comme ci-dessus

$$\lim \left[\frac{J_m(r y) J_m(r_1 y)}{(J'_0(y))^2} \right]^{y=0} = 0$$

pour $m > 1$. Mais comme, $J'_0(y) = -J_1(y)$, on aura,
pour $m = 1$,

$$\begin{aligned} & \lim \left[\frac{J_1(r y) J_1(r_1 y)}{(J'_0(y))^2} \right]^{y=0} = r r_1 \\ & = -2 \sum \frac{r J_0(k' r) J_1(k' r_1) - r_1 J_1(k' r) J_2(k' r_1)}{k' (J_0(k'))^2}, \end{aligned}$$

et, comme le dernier terme s'évanouit pour $r_1 < r$,

$$-r_1 = 2 \sum \frac{J_0(k' r) J_1(k' r_1)}{k' (J_0(k'))^2} \text{ pour } r_1 < r.$$

On doit ensuite chercher la valeur de

$$\frac{J_0(y r) J_0(y r_1)}{(J'_0(y))^2}.$$

Ici, l'on peut décomposer la fraction comme à l'ordinaire,
mais on doit tenir compte de la racine $y = 0$ de l'équa-
tion $J'_0(y) = 0$, valeur qui ne satisfait pas à l'équation
 $J_0(y r) J_0(y r_1) = 0$.

On obtiendra ainsi

$$\frac{J_0(yr) J_0(yr_1)}{(J'_0(y))^2} = \frac{4}{y^2} + 2 \sum \frac{1}{J_0(k')} \frac{d}{dk'} \left(\frac{J_0(k'r) J_0(k'r_1)}{(k' - y) J''_0(k')} \right),$$

ou bien, si l'on fait tendre y vers zéro

$$\begin{aligned} & \lim \left(\frac{4}{y^2} - \frac{J_0(yr) J_0(yr_1)}{(J'_0(y))^2} \right)^{y=0} \\ &= -2 \sum \frac{r J_1(k'r) J_0(k'r_1) + r_1 J_1(k'r_1) J_0(k'r)}{k' (J_0(k'))^2}. \end{aligned}$$

Mais

$$\lim \left(\frac{4}{y^2} - \frac{J_0(yr) J_0(yr_1)}{(J'_0(y))^2} \right)^{y=0} = r^2 + r_1^2 - 1.$$

Si l'on suppose $r < r_1$, on aura donc

$$-2 \sum \frac{r_1 J_1(k'r_1) J_0(k'r)}{k' (J_0(k'))^2} = r_1^2 - 1.$$

A présent nous pouvons sans difficulté trouver la valeur de la série

$$u = 2 \sum \frac{J_0(k'r)}{(J_0(k'))^2} \int_0^1 J_0(k'r') r' f(r') dr'.$$

Si r_1 et r_2 sont infiniment rapprochés, et sous la condition que $f(r)$ reste finie et continue, on aura

$$\begin{aligned} \int_{r_1}^{r_2} J_0(k'r') r' f(r') dr' &= f(r_1) \int_{r_1}^{r_2} J_0(k'r') r' dr' \\ &= \frac{f(r_1)}{k'} [r_2 J_1(k'r_2) - r_1 J_1(k'r_1)], \end{aligned}$$

et, par conséquent,

$$\begin{aligned} & 2 \sum \frac{J_0(k'r)}{(J_0(k'))^2} \int_{r_1}^{r_2} J_0(k'r') r' f(r') dr' \\ &= f(r_1) \left[2 \sum \frac{r_2 J_0(k'r) J_1(k'r_2)}{k' (J_0(k'))^2} - 2 \sum \frac{r_1 J_0(k'r) J_1(k'r_1)}{k' (J_0(k'))^2} \right] \\ &= f(r_1) [-r_2^2 + r_1^2], \quad \text{si } r > r_2 \text{ ou } r < r_1. \end{aligned}$$

Si $r_2 > r > r_1$, l'expression est égale à $(1 - r_2^2 + r_1^2)f(r)$.
 Si r_1 et r_2 sont infiniment rapprochés, on peut poser $r_2 - r_1 = dr_1$, $r_1 + r_2 = 2r_1$, et en prenant la somme de tous les éléments de l'intégrale, on obtiendra

$$u = f(r) - 2 \int_0^r f(r_1) dr_1,$$

ce qui démontre la proposition de Lorenz.

SUR LES NOMBRES PREMIERS.

SUR LES NOMBRES PREMIERS.

TIDSSKRIFT FOR MATEMATIK 1878, P. 1—3.

Il est toujours possible d'obtenir une expression exacte de la forme

$$A_{p,n} = B_p \cdot n + C_{p,n} \quad (1)$$

du nombre $A_{p,n}$ des termes de la suite 1, 2, 3 ... n qui ne sont pas divisibles par les nombres premiers 2, 3 ... p , $C_{p,n}$, étant une fonction périodique qui reprend les mêmes valeurs quand n augmente du produit des nombres premiers donnés.

Cette fonction périodique peut d'ailleurs, grâce à des méthodes bien connues, être déterminée au moyen des fonctions trigonométriques. En effectuant ce calcul, j'ai trouvé que le résultat peut être exprimé sous forme simple par la fonction périodique

$$n_m = \frac{\sin \frac{n\pi}{m}}{\sin \frac{\pi}{m}} - \frac{\sin \frac{2n\pi}{m}}{\sin \frac{2\pi}{m}} + \dots - \frac{\sin \frac{(m-1)n\pi}{m}}{\sin \frac{(m-1)\pi}{m}}, \quad (2)$$

où m est un nombre impair, n un nombre pair. On obtiendra ce résultat de la manière la plus facile, si l'on exécute la sommation, les fonctions trigonométriques étant exprimées en exponentielles imaginaires. On trouvera de cette manière

* NOTE 1.

$$n_m = n - 2mr^* \quad (3)$$

r étant le nombre des termes de la série 1, 3, 5 ... $(n-1)$ divisibles par m . Il s'ensuit que

$$r = \frac{n-n_m}{2m}; \quad (4)$$

ainsi r est exprimé par la fonction trigonométrique n_m .

Dans la série 1, 2, 3 ... n (n pair) se trouvent $\frac{1}{2}n$ nombres qui ne sont pas divisibles par 2. Des nombres restants impairs $\frac{n-n_s}{2 \cdot 3}$ sont d'après (4) divisibles par 3. Le nombre des termes de la série 1, 2, 3 ... n qui ne sont divisibles ni par 2, ni par 3, est donc $\frac{1}{2}(n - \frac{1}{3}(n-n_s))$. Si l'on en soustrait, de plus, le nombre des nombres divisibles par 5, on soustrait deux fois les nombres divisibles par 3.5; c'est pourquoi l'on doit ajouter le dernier nombre à la différence. Par conséquent, le nombre des nombres non divisibles par 2, 3 et 5 est

$$\frac{1}{2} \left(n - \frac{1}{3}(n-n_s) - \frac{1}{5}(n-n_s) + \frac{1}{3 \cdot 5}(n-n_{s,s}) \right).$$

On reconnaît de cette manière sans difficulté qu'on a dans la formule générale (1)

$$B_p = \frac{1}{2} \left(1 - \frac{1}{3} \right) \left(1 - \frac{1}{5} \right) \dots \left(1 - \frac{1}{p} \right), \quad (5)$$

et, en se servant d'une notation symbolique,

$$C_{p,n} = \frac{1}{2} \left[1 - \left(1 - \frac{n_s}{3} \right) \left(1 - \frac{n_s}{5} \right) \dots \left(1 - \frac{n_p}{p} \right) \right], \quad (6)$$

où l'on doit, après avoir exécuté la multiplication, remplacer $n_m n_q$ par n_{mq} . Ces formules sont valables pour

n pair; mais, comme n est lui-même un nombre non premier, on doit avoir

$$A_{p,n-1} = A_{p,n}. \quad (7)$$

Si, dans l'expression (2), l'on fait parcourir à n toutes les valeurs paires à partir d'un nombre arbitraire n_1 jusqu'à $n_1 + 2m - 2$, et si l'on fait la sommation des valeurs correspondantes n_m , le résultat s'évanouira. Il s'ensuit que pareillement la somme des $C_{p,n}$ s'évanouira, si n parcourt les valeurs d'une série de nombres pairs consécutifs dont le nombre est $3 \cdot 5 \dots p$.

D'après (7) on peut poser

$$A_{p,n-1} = B_p(n-1) + B_p + C_{p,n},$$

et, par conséquent, on aura, pour n pair,

$$C_{p,n-1} = B_p + C_{p,n}.$$

La valeur moyenne de $C_{p,n}$ est, par conséquent, si n parcourt toute la série des nombres pairs et impairs jusqu'à $2 \cdot 3 \cdot 5 \dots p$, égale à $\frac{1}{2} B_p$.

Si $\varphi(n)$ désigne le nombre des nombres premiers dans la suite des nombres entiers $1, 2, 3 \dots n$, et si p est le plus grand nombre premier inférieur à $\sqrt{n}$, on aura

$$A_{p,n} = \varphi(n) - \varphi(p) + 1, \quad (8)$$

car $A_{p,n}$ est le nombre des nombres premiers inférieurs à n , à l'exception des nombres premiers eux mêmes, $2, 3, 5 \dots p$ dont le nombre est $\varphi(p) - 1$. On aura donc approximativement, en négligeant la partie périodique,

$$\varphi(n) - \varphi(p) + 1 = B_p \cdot n,$$

expression assez bien connue. La fonction $\varphi(n)$ a été l'objet de recherches si approfondies que je n'ai rien

d'essentiellement nouveau à en dire. Pourtant, je ne me souviens pas d'avoir vu aucun essai fait en vue de déterminer $\varphi(n)$ par une équation différentielle, problème qui pourtant peut facilement être résolu.

Qu'on pose approximativement

$$\varphi(p^2) - \varphi(p) = B_p \cdot p^2,$$

et, p_1 étant le nombre premier qui suit p ,

$$\varphi(p_1^2) - \varphi(p_1) = B_{p_1} \cdot p_1^2 = B_p \left(1 - \frac{1}{p}\right) p_1^2.$$

Si l'on considère $\varphi(p)$ comme une fonction continue d'une variable continue, et si l'on développe cette fonction en série de Taylor, on obtiendra, dans l'hypothèse où $p_1 - p$ est très petit en comparaison de p ,

$$\frac{d}{dp}(\varphi(p^2) - \varphi(p))(p_1 - p) = B_p \cdot 2p(p_1 - p) - B_p p.$$

De la même manière on déduira de l'équation

* NOTE 2. $\varphi(p_1) - \varphi(p) = 1$ l'équation $\varphi'(p)(p_1 - p) = 1^*$; puis, par élimination de $p_1 - p$ et de B_p , on obtiendra

$$\frac{d}{dp} l \left(\varphi(p^2) - \varphi(p) \right) = \frac{2 - \varphi'(p)}{p}. \quad (9)$$

Pour les très grandes valeurs de p on déduira facilement de cette équation la formule de Gauss $\varphi(p)$

$= \int_0^p \frac{dx}{lx}$, la première approximation étant

$$\varphi(p^2) - \varphi(p) = \int_p^{p^2} \frac{dx}{lx} = \frac{p^2}{2lp},$$

valeur qui satisfait à l'équation différentielle.* D'ail- * NOTE 3.
 leurs, l'intégrale complète de l'équation (9) doit contenir
 pour les valeurs finies de p une fonction arbitraire ou
 un nombre illimité de constantes arbitraires.

Mais $\cos \frac{q}{2}$ est différent de zéro, si $(2t+1)$ n'est pas un multiple de m . Dans le cas contraire $\frac{2t+1}{m}$ est un nombre entier et impair, m étant un tel nombre, et par suite $\cos q = -1$,

$$s = -2(m-1).$$

Il en résulte que

$$n_m = 2 \left(\frac{n}{2} - r \right) - 2(m-1)r = n - 2mr.$$

NOTE 2. Tout le procédé de Lorenz repose sur une hypothèse qui n'est pas démontrée et ne peut donner qu'un résultat incertain, puisqu'on ne connaît pas les limites des termes négligés.

De l'équation

$$\varphi(p_1) - \varphi(p) = 1$$

Lorenz déduit, par exemple, au moyen de la série de Taylor, l'équation

$$\varphi'(p)(p_1 - p) = 1,$$

en négligeant les termes d'ordre supérieur par rapport à $p_1 - p$. Mais on ne sait pas si les termes négligés ne sont pas très grands en comparaison du terme conservé 1.

NOTE 3. Si l'on pose $\varphi(p) = \int_0^p \frac{dx}{lx}$ on aura

$$\varphi(p^2) - \varphi(p) = \int_p^{p^2} \frac{dx}{lx} = \frac{p^2}{2lp} - \frac{p}{lp} - \int_p^{p^2} \frac{dx}{(lx)^2}.$$

En négligeant les deux derniers termes, qui sont petits en comparaison du terme $\frac{p^2}{2l(p)}$, on obtiendra l'expression du texte. On reconnaît facilement que l'équation différentielle est satisfaite, si l'on pose en même temps

$$\varphi(p^2) - \varphi(p) = \frac{p^2}{2l(p)}$$

et

$$\varphi'(p) = \frac{1}{l(p)}.$$

RECHERCHES ANALYTIQUES

SUR

LES NOMBRES DE NOMBRES PREMIERS.

RECHERCHES ANALYTIQUES SUR LES NOMBRES DE NOMBRES PREMIERS.

VIDENSKABERNES SELSKABS SKR. V (SÉRIE 6), 1891, P. 427—450.

Dans le mémoire célèbre „Über die Anzahl der Primzahlen unter einer gegebenen Grösse“*, Riemann s'est servi comme point de départ de l'équation d'Euler**

$$\sum n^r \prod (1 - p^r) = 1, \quad (1)$$

où la somme $\sum$ porte sur tous les nombres entiers n , et le produit $\prod$ sur tous les nombres premiers.

De ce point de départ l'analyse ne mène pas à la détermination directe de la fonction $\theta(x)$, qui exprime le nombre des nombres premiers inférieurs à x (x inclus), mais plutôt à une autre fonction $\vartheta(x)$, que M. le Dr. J.-P. Gram a appelée *le nombre des puissances divisées des nombres premiers****, et qui est définie par l'équation

$$\vartheta(x) = \theta(x) + \frac{1}{2}\theta(x^{\frac{1}{2}}) + \frac{1}{3}\theta(x^{\frac{1}{3}}) + \dots \quad (2)$$

C'est cette fonction que Riemann a cherché à déterminer sous forme analytique. Riemann a de plus démontré qu'on peut, en partant de cette fonction, retrouver la fonction $\theta(x)$; mais il faut pourtant remarquer que toujours, pour représenter analytiquement $\vartheta(x)$, on devra,

* Monatsberichte d. K. Akad. d. W. zu Berlin 3 nov. 1859, p. 671.

** Euler: Introductio in Analysin Infinitorum, t. 1, p. 237, 1748.

*** Videnskabernes Selsk. Skr. II (série 6), p. 8.

selon toute apparence, contrairement à la détermination théoriquement exacte, partir de la limite (infinie) supérieure de la série des nombres et puis se servir d'un calcul approché pour les valeurs de x extrêmement grandes, ce qui exclut la détermination de $\vartheta(x)$ pour les nombres petits et aussi la détermination de $\theta(x)$ au moyen de $\vartheta(x)$. Cette dernière fonction peut toutefois, aussi bien que le nombre même des nombres premiers, être représentée au moyen de tableaux que l'on forme en se fondant sur les dénombrements de nombres premiers; c'est pourquoi il suffit, en fait, de déterminer $\vartheta(x)$.

Bien qu'on ne puisse pas accepter la démonstration de la formule de $\vartheta(x)$, telle que Riemann l'a présentée, le travail de Riemann a fait ressortir ce résultat final, que l'expression de la partie apériodique de $\vartheta(x)$, le logarithme intégral $\text{Li}(x)$, est en pratique parfaitement suffisante, au moins entre les limites entre lesquelles est connu jusqu'ici le nombre des nombres premiers, et de beaucoup supérieure à toutes les formules développées antérieurement à la suite de recherches empiriques.

C'est ce fait que Glaisher* principalement a établi pour chaque intervalle de 50000 nombres jusqu'à 9 millions et pour 10 et 100 millions d'après les nombres de nombres premiers calculés par Meissel**. Les écarts sont représentés graphiquement et, pour la formule de Riemann, le diagramme indique une répartition uniforme des écarts négatif et des écarts positifs. *En ce*

* James Glaisher: Factor tables for the sixth Million. London 1883.

** Meissel a plus tard (Math. Ann., tome 15, p. 251) calculé le nombre des nombres premiers jusqu'à 1000 millions. Son résultat (50847478) ne dépasse que de 23 unités le résultat de Gram, calculé au moyen de la formule de Riemann.

qui concerne les nombres plus grands, ces écarts semblent se développer successivement en de longues périodes, de plus en plus régulières.

D'après cela, la marche naturelle pour la recherche analytique des nombres de nombres premiers sera de choisir le même point de départ que Riemann, à savoir la formule d'Euler (1), parce qu'elle mène immédiatement à la détermination de $\vartheta(x)$, dans laquelle la forme de la partie apériodique s'est montrée simple en pratique; de chercher à déterminer analytiquement cette partie apériodique avec une exactitude aussi grande que possible; puis de chercher à déterminer la partie périodique de $\vartheta(x)$ par un développement en série, de manière à faire apparaître en premier lieu les termes qui sont de l'ordre le plus élevé par rapport à x , pour faire ressortir analytiquement de cette manière, s'il est possible, les longues périodes mentionnées ci-dessus qui apparaissent pour les nombres très grands. C'est cette marche que j'ai suivie dans mes recherches.

Nous commencerons par une série de la forme

$$(2^r + 3^r + 4^r \dots)^s = \alpha^s(2) 2^r + \alpha^s(3) 3^r + \dots + \alpha^s(x) x^r \dots (3),$$

où r est une quantité arbitraire, s un nombre entier et $\alpha^s(x)$ un coefficient qui indique le nombre des manières différentes suivant lesquelles x peut être formé par multiplication de s nombres entiers, le nombre 1 non compris.

Si nous posons de plus

$$A^s(x) = \alpha^s(2) + \alpha^s(3) + \dots + \alpha^s(x), \quad (4)$$

$\alpha^s(x)$ étant le dernier terme de la série (3), on aura

$$A^s(x) - A^s(x-1) = \alpha^s(x). \quad (5)$$

Si l'on prend le logarithme des deux membres de l'équation identique (1), et si l'on emploie le développement en série de $\log(1+y)$ suivant les puissances croissantes de y , sans tenir compte de la convergence de la série, on obtiendra de nouveau des développements identiques* qui, pour $r = 0$, font ressortir l'équation bien connue

$$\vartheta(x) = \frac{A^1(x)}{1} - \frac{A^2(x)}{2} + \frac{A^3(x)}{3} \dots \quad (6)$$

Cette série est finie; car tous les premiers coefficients α^s jusqu'à $\alpha^s(2^s)$ s'évanouissent, et par conséquent $A^s(x)$ s'évanouira d'après (4) pour $x < 2^s$ ou $s > \frac{\log x}{\log 2}$. Nous remplacerons les quantités $A^s(x)$ par d'autres quantités $B^s(x)$, qui se prêtent mieux au calcul suivant; pour les définir, nous modifierons légèrement la série (3) en posant

$$(\frac{1}{2} + 2^r + 3^r + 4^r \dots)^s = \beta^s(1) + \beta^s(2) 2^r + \beta^s(3) 3^r + \dots + \beta^s(x) x^r + \dots \quad (7)$$

et

$$B^s(x) = \beta^s(1) + \beta^s(2) + \dots + \beta^s(x). \quad (8)$$

Les coefficients $\beta^s(x)$ seront de même les expressions des nombres de manières différentes suivant lesquelles x peut être formé par multiplication de s nombres entiers: mais à présent le nombre 1 est compris parmi les s facteurs, seulement on ne prend que la moitié du nombre des décompositions où figure le facteur 1.

* Cfr. les remarques de J. Petersen sur le mémoire de Gram. Vid. Selsk. Overs. 1884, p. 14.

On aura alors

$$A^s(x) = B^s(x) - \frac{s}{2} B^{s-1}(x) + \frac{s(s-1)}{2 \cdot 4} B^{s-2}(x) + \dots \\ + (-1)^{s-1} \frac{s}{2^{s-1}} B^1(x) + \left(-\frac{1}{2}\right)^s, * \quad (9) * \text{NOTE 1.}$$

et on peut remplacer (6) par une série de la forme

$$\vartheta(x) = -a_0 + \frac{a_1 B^1(x)}{1} - \frac{a_2 B^2(x)}{2} + \dots \pm a_{s_1} \frac{B^{s_1}(x)}{s_1}, \quad s_1 > \frac{\log x}{\log 2} - 1, \quad (10)$$

où les coefficients a sont déterminés par

$$a_p = 1 + \frac{p}{2} + \frac{p(p+1)}{2 \cdot 4} + \dots + \frac{p(p+1) \dots (s_1-1)}{2 \cdot 4 \dots (2s_1-2p)}, \quad p > 0, \quad (11)$$

$$a_0 = \frac{1}{1} \cdot \frac{1}{2} + \frac{1}{2} \cdot \frac{1}{2^2} + \dots + \frac{1}{s_1} \cdot \frac{1}{2^{s_1}}. \quad (12)$$

On peut remarquer que le nombre s_1 peut être choisi arbitrairement pourvu, qu'il dépasse $\frac{\log x}{\log 2}$ et peut être pris égal à ∞ , valeur à laquelle correspond $a_p = 2^p$, $a_0 = \log 2$.

Mais si l'on ne considère séparément qu'une partie des fonctions $\vartheta(x)$ et $B^s(x)$, il va sans dire qu'il n'est permis de poser $x = \infty$ que si la série correspondante pour cette valeur de s_1 devient convergente.

Comme fondement de mes recherches j'ai choisi la formule de sommation de Poisson, d'après laquelle on a exactement

$$\frac{1}{2} + 2^r + 3^r + \dots + x^r = \int_{\epsilon_1}^{x+\frac{1}{2}} dx_1 x_1^r (1 + 2 \sum \cos 2\pi m_1 x_1), \quad m_1 = 1, 2, \dots \infty. \quad (13)$$

Si les limites sont choisies d'une manière convenable, on peut en outre exprimer la puissance s -ième de cette

série par un intégrale s -uple

$$\int dx_1 x_1^r (1 + 2 \sum \cos 2\pi m_1 x_1) \int dx_2 x_2^r (1 + 2 \sum \cos 2\pi m_2 x_2) \dots \\ \int dx_s x_s^r (1 + 2 \sum \cos 2\pi m_s x_s),$$

où $m_1, m_2, \dots, m_s$ parcourent chacun séparément les valeurs $1, 2 \dots \infty$.

Ici s'introduisent les variables nouvelles

$$u_1 = x_1 x_2 \dots x_s, \quad u_2 = x_2 x_3 \dots x_s, \quad u_{s-1} = x_{s-1} x_s, \quad u_s = x_s.$$

Comme la limite inférieure de toutes les variables x est l'unité, ce sera aussi la limite inférieure des variables nouvelles. La limite supérieure de $u_s = \frac{u_{s-1}}{x_{s-1}}$ est u_{s-1} ; de même, celle de u_{s-1} sera u_{s-2} et ainsi de suite, jusqu'à u_1 , dont la limite supérieure sera désignée par u . De plus, nous aurons

$$dx_s = du_s, \quad dx_{s-1} = \frac{du_{s-1}}{u_s}, \quad dx_{s-2} = \frac{du_{s-2}}{u_{s-1}} \dots dx_1 = \frac{du_1}{u_2}.$$

De cette manière on peut représenter l'expression (7) par une intégrale s -uple, et l'on verra qu'on peut arrêter le développement à $B^s(x)x^r$, pris comme dernier

*NOTE 2. terme, en posant la limite $u = x + \frac{1}{2}$. *

Si l'on fait ensuite $r = 0$, on obtiendra

$$B^s(x) = \int_1^u du_1 \int_1^{u_1} \frac{du_2}{u_2} \dots \int_1^{u_{s-1}} \frac{du_s}{u_s} \left(1 + 2 \sum \cos \mu_1 \frac{u_1}{u_2} \right) \dots \\ \left(1 + 2 \sum \cos \mu_{s-1} \frac{u_{s-1}}{u_s} \right) (1 + 2 \sum \cos \mu_s u_s), \quad (14)$$

où l'on a posé, pour abrégé,

$$2\pi m_1 = \mu_1, \quad 2\pi m_2 = \mu_2, \quad \dots, \quad 2\pi m_s = \mu_s.$$

De même que tous les éléments de l'intégrale simple (13) s'évanouissent, si x_1 n'est pas un nombre entier, de même tous les éléments de cette intégrale multiple s'évanouiront, à moins que toutes les fractions

$$\frac{u_1}{u_2}, \frac{u_2}{u_3} \dots \frac{u_{s-1}}{u_s}, \frac{u_s}{1}$$

soient en même temps des nombres entiers.

Si l'on effectue la multiplication de tous les facteurs entre parenthèses, les termes du produit pourront être ordonnés d'après le nombre des signes de sommation qui y entrent et l'on peut démontrer que toutes les intégrales qui contiennent le même nombre de signes de sommation sont égales. En effet, si f, g, h désignent des fonctions arbitraires, on peut démontrer l'identité

$$\begin{aligned} & \int_1^{u_{p-1}} \frac{du_p}{u_p} f\left(\frac{u_{p-1}}{u_p}\right) \int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} g\left(\frac{u_p}{u_{p+1}}\right) h(u_{p+1}) \\ &= \int_1^{u_{p-1}} \frac{du_p}{u_p} g\left(\frac{u_{p-1}}{u_p}\right) \int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} f\left(\frac{u_p}{u_{p+1}}\right) h(u_{p+1}), \quad (15) \end{aligned}$$

dont les deux membres ne diffèrent que par la permutation de f avec g . Le premier se transforme en

$$\int_1^{u_{p-1}} \frac{du_p}{u_p} f\left(\frac{u_{p-1}}{u_p}\right) \int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} g(u_{p+1}) h\left(\frac{u_p}{u_{p+1}}\right)$$

si l'on remplace u_{p+1} par $\frac{u_p}{u_{p+1}}$. Si l'on introduit la notation

$$\phi(u, u_{p+1}) = \int \frac{du}{u} f\left(\frac{u_{p-1}}{u}\right) h\left(\frac{u}{u_{p+1}}\right)$$

l'intégrale ci-dessus peut être exprimée par

$$\begin{aligned} \int_1^{u_{p-1}} \left[\frac{d}{du_p} \left(\int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} g(u_{p+1}) \phi(u_p, u_{p+1}) \right) - \frac{1}{u_p} g(u_p) \phi(u_p, u_p) \right] \\ = \int_1^{u_{p-1}} \frac{du_{p+1}}{u_{p+1}} g(u_{p+1}) \phi(u_{p-1}, u_{p+1}) - \int_1^{u_{p-1}} \frac{du_p}{u_p} g(u_p) \phi(u_p, u_p), \end{aligned}$$

où u_{p+1} peut dans la première intégrale être remplacé par u_p . Les deux intégrales peuvent ensuite être réunies en une seule

$$\begin{aligned} \int_1^{u_{p-1}} \frac{du_p}{u_p} g(u_p) (\phi(u_{p-1}, u_p) - \phi(u_p, u_p)) \\ = \int_1^{u_{p-1}} \frac{du_p}{u_p} g(u_p) \int_{u_p}^{u_{p-1}} \frac{du_{p+1}}{u_{p+1}} f\left(\frac{u_{p-1}}{u_{p+1}}\right) h\left(\frac{u_{p+1}}{u_p}\right), \end{aligned}$$

et si l'on y remplace u_p par $\frac{u_{p-1}}{u_p}$, puis u_{p+1} par $\frac{u_{p-1}}{u_p} u_{p+1}$, l'expression sera identique au second membre de l'identité (15), qui se trouve par là démontrée.

Si nous revenons maintenant au produit en question et si nous considérons séparément l'intégrale qui contient $\sum \cos \mu_p \frac{u_p}{u_{p+1}}$, mais non $\sum \cos \mu_{p-1} \frac{u_{p-1}}{u_p}$, nous pourrions poser

$$\begin{aligned} \dots \int_1^{u_{p-1}} \frac{du_p}{u_p} \int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} \sum \cos \mu_p \frac{u_p}{u_{p+1}} \dots \\ = \dots \int_1^{u_{p-1}} \frac{du_p}{u_p} \sum \cos \mu_{p-1} \frac{u_{p-1}}{u_p} \int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} \dots \end{aligned}$$

De cette manière on peut transporter toutes les sommes Σ de droite à gauche en faisant un changement correspondant des indices, par où toutes les intégrales pour lesquelles le nombre des sommes est le même deviendront identiques.*

* NOTE 3.

Les intégrales qui contiennent p signes de sommation prendront donc la forme

$$\int_1^u du_1 \int_1^{u_1} \frac{du_2}{u_2} \Sigma \cos \mu_1 \frac{u_1}{u_2} \dots \int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} \Sigma \cos \mu_p \frac{u_p}{u_{p+1}} \int_1^{u_{p+1}} \frac{du_{p+2}}{u_{p+2}} \dots \int_1^{u_{s-1}} \frac{du_s}{u_s}.$$

Si l'on désigne cette expression par X_p^s et si l'on calcule les dernières intégrales, on obtiendra, pour $p < s$,

$$X_p^s = \int_1^u du_1 \int_1^{u_1} \frac{du_2}{u_2} \Sigma \cos \mu_1 \frac{u_1}{u_2} \dots \int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} \Sigma \cos \mu_p \frac{u_p}{u_{p+1}} \cdot \frac{(\log u_{p+1})^{s-p-1}}{(s-p-1)!}. \quad (16)$$

Si l'on a $p = s$, on pose $\int_1^u du_1 = u \int_1^u \frac{du_1}{u_1} \cdot \frac{u_1}{u}$; puis on peut permuter $\frac{u_1}{u}$ et le signe de sommation suivant en changeant les indices et ainsi de suite. La dernière intégrale deviendra

$$\int_1^{u_{s-1}} \frac{du_s}{u_{s-1}} \Sigma \cos \mu_s u_s = \Sigma \frac{\sin \mu_s u_{s-1}}{\mu_s u_{s-1}},$$

et l'expression totale prendra la forme

$$X_s^s = u \int_1^u \frac{du_1}{u_1} \Sigma \cos \mu_1 \frac{u}{u_1} \dots \int_1^{u_{s-2}} \frac{du_{s-1}}{u_{s-1}} \Sigma \cos \mu_{s-1} \frac{u_{s-2}}{u_{s-1}} \cdot \Sigma \frac{\sin \mu_s u_{s-1}}{\mu_s u_{s-1}}. \quad (17)$$

Si l'on porte ces valeurs de X_p^s et X_s^s dans l'équation (14), elle prendra la forme

$$B^s(x) = X_0^s + \frac{s}{1} X_1^s + \frac{s(s-1)}{1 \cdot 2} X_2^s + \dots + \frac{s}{1} X_{s-1}^s + X_s^s. \quad (18)$$

Ici se présente tout d'abord le problème de déterminer la partie apériodique de $B^s(x)$, que je désignerai par $\bar{B}^s(x)$ en employant, comme dans tout ce qui suit, un trait horizontal au dessus du signe fonctionnel pour représenter la partie apériodique de la fonction.

La dernière intégrale qui figure dans (16) est

$$\begin{aligned} & \int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} 2 \sum \cos \mu_p \frac{u_p}{u_{p+1}} \cdot \frac{(\log u_{p+1})^{s-p-1}}{(s-p-1)!} \\ &= \int_1^{u_p} \frac{du_{p+1}}{u_{p+1}} 2 \sum \cos \mu_p u_{p+1} \frac{(\log u_p - \log u_{p+1})^{s-p-1}}{(s-p-1)!}. \end{aligned}$$

Si l'on décompose la dernière intégrale en les deux suivantes

$$\begin{aligned} & \int_1^\infty \frac{du_{p+1}}{u_{p+1}} 2 \sum \cos \mu_p u_{p+1} \frac{(\log u_p - \log u_{p+1})^{s-p-1}}{(s-p-1)!} \\ & - \int_{u_p}^\infty \frac{du_{p+1}}{u_{p+1}} 2 \sum \cos \mu_p u_{p+1} \frac{(\log u_p - \log u_{p+1})^{s-p-1}}{(s-p-1)!}, \end{aligned}$$

la première peut être considérée comme la partie apériodique, la seconde comme la partie périodique de l'intégrale. La première peut être mise sous la forme

$$C_0 \frac{(\log u_p)^{s-p-1}}{(s-p-1)!} + C_1 \frac{(\log u_p)^{s-p-2}}{(s-p-2)!} + C_2 \frac{(\log u_p)^{s-p-3}}{(s-p-3)!} + \dots,$$

où les constantes C_n sont définies par l'équation

$$C_n = \frac{(-1)^n}{n!} \int_1^\infty \frac{du}{u} (\log u)^n 2 \sum' \cos \mu u; * \quad (19) \quad * \text{ NOTE 4.}$$

on trouvera ainsi

$$\begin{aligned} C_0 &= 0,07721 \dots \\ C_1 &= 0,07281 \dots \\ C_2 &= -0,00484 \dots \\ C_3 &= -0,00034 \dots \\ C_4 &= 0,00009 \dots \end{aligned}$$

Si de plus, l'on pose, pour abrégé,

$$C_0 \frac{\partial}{\partial (\log u)} + C_1 \frac{\partial^2}{(\partial \log u)^2} + C_2 \frac{\partial^3}{(\partial \log u)^3} + \dots = \Delta_{\log u}, \quad (20)$$

la partie apériodique de la dernière intégrale de (16) prend la forme

$$\Delta_{\log u_p} \frac{(\log u_p)^{s-p}}{(s-p)!}.$$

Cette expression étant portée dans (16), la dernière intégrale peut s'écrire

$$\begin{aligned} & \int_1^{u_{p-1}} \frac{du_p}{u_p} 2 \sum' \cos \mu_{p-1} \frac{u_{p-1}}{u_p} \cdot \Delta_{\log u_p} \frac{(\log u_p)^{s-p}}{(s-p)!} \\ &= \int_1^{u_{p-1}} \frac{du_p}{u_p} 2 \sum' \cos \mu_{p-1} u_p \cdot \Delta_{\log u_{p-1}} \frac{(\log u_{p-1} - \log u_p)^{s-p}}{(s-p)!}, \end{aligned}$$

d'où l'on conclut comme ci-dessus que la partie apériodique est déterminée par

$$\Delta_{\log u_{p-1}} \Delta_{\log u_{p-1}} \frac{(\log u_{p-1})^{s-p+1}}{(s-p+1)!} = \Delta_{\log u_{p-1}}^2 \frac{(\log u_{p-1})^{s-p+1}}{(s-p+1)!}.$$

De cette manière on reconnaît facilement que la partie apériodique de (16) peut être exprimée par

$$\bar{X}_p^s = \int_1^u d u_1 \Delta_{\log u_1}^p \frac{(\log u_1)^{s-1}}{(s-1)!}. \quad (21)$$

En conséquence, la partie apériodique totale de (18) peut se mettre sous les formes symboliques

$$\bar{B}_{(x)}^s = \int_1^u d u_1 (1 + \Delta_{\log u_1})^s \frac{(\log u_1)^{s-1}}{(s-1)!} = \int_0^{\log u} d v e^v (1 + \Delta_v)^s \frac{v^{s-1}}{(s-1)!}. \quad (22)$$

Tant que nous nous trouvons entre les limites de x qu'on a atteintes jusqu'ici par le dénombrement ou le calcul exact du nombre des nombres premiers, limites qui correspondent respectivement à $\log x = 16$ et $\log x = 21$, les difficultés pour la détermination de la partie apériodique de $\vartheta(x)$, si l'on porte les valeurs trouvées de $\bar{B}_{(x)}^s$ dans l'équation (10), ne seront pas insurmontables, s_1 étant le plus grand nombre entier inférieur à $\frac{\log x}{\log 2}$. Le résultat ne peut guère être mis sous une forme notablement plus simple, si l'on veut tenir compte de toutes les constantes C_n qui entrent dans Δ_v ; mais si l'on ne tient compte que de la première C_0 , la série sera convergente pour $s_1 = \infty$, et la sommation pourra facilement être effectuée. On aura évidemment pour $s_1 = \infty$, $C_1 = 0$, $C_2 = 0 \dots$

$$\begin{aligned} \bar{\vartheta}(x) &= -\log 2 + \sum_{s=1}^{s=\infty} (-1)^{s-1} 2^s \int_0^{\log u} d v e^v \left(1 + C_0 \frac{d}{dv}\right)^s \frac{v^{s-1}}{s!} \\ &= -\log 2 + \int_0^{\log u} d v e^v \left[\frac{1}{v} - \frac{e^{-2v}}{v} + \frac{2C_0}{1!} \frac{d}{dv} \left(v \cdot \frac{e^{-2v}}{v} \right) - \frac{2^2 C_0^2}{2!} \frac{d^2}{dv^2} \left(v^2 \frac{e^{-2v}}{v} \right) + \dots \right], \end{aligned}$$

et, si l'on effectue la sommation au moyen de la série

* NOTE 5. de Lagrange*, on aura

$$\left. \begin{aligned} \bar{\vartheta}(x) &= -\log 2 + \int_0^{\log u} dv e^v \left(\frac{1}{v} - \frac{e^{-\frac{2v}{1+2C_0}}}{v} \right) \\ &= -\log 2 + Li(u) - Li\left(u^{-\frac{1-2C_0}{1+2C_0}}\right). \end{aligned} \right\} \quad (23)$$

Ici on a $\frac{1-2C_0}{1+2C_0} = 0,73245 \dots$, ce qui fait ressortir que l'expression trouvée, pour les valeurs de x très grandes, ne diffère pas essentiellement du logarithme intégral de x et qu'on pourrait même poser $C_0 = 0$, ce qui n'aurait pas sensiblement modifié la valeur de $\bar{\vartheta}(x)$. Dans ce cas, le résultat prendra la forme simple

$$\bar{\vartheta}(x) = -\log 2 + \int_{-\log u}^{+\log u} \frac{dv}{v} e^v.$$

En pratique ces résultats sont donc en bonne concordance avec ceux de Riemann.

La partie périodique de X_p^s , $p < s$, n'a que peu d'importance en comparaison de X_s^s , et, de plus, comme la première peut être déduite de la dernière, je me bornerai dans ce qui suit à la discussion de la seule fonction X_s^s qui est définie par l'équation (17).

Nous commencerons par considérer le cas $s = 2$, à savoir

$$X_2^2 = u \int_1^u \frac{du_1}{u_1} 2 \sum' \cos \mu_1 \frac{u}{u_1} \cdot 2 \sum' \frac{\sin \mu_2 u_1}{\mu_2 u_1}.$$

L'intégration par parties transforme X_2^2 en

$$\begin{aligned} X_2^2 &= \int_1^u du_1 2 \sum' \frac{\sin \mu_1 \frac{u}{u_1}}{\mu_1} \cdot 2 \sum' \cos \mu_2 u_1 \\ &= 2 \sum' \sum' \int_1^u \frac{du_1}{\mu_1} \left(\sin \left(\mu_1 \frac{u}{u_1} + \mu_2 u_1 \right) + \sin \left(\mu_1 \frac{u}{u_1} - \mu_2 u_1 \right) \right). \end{aligned}$$

On pose ici

$$\begin{aligned}\mu_1 \frac{u}{u_1} + \mu_2 u_1 &= 2v \sqrt{\mu_1 \mu_2 u}, & \mu_1 \frac{u}{u_1} - \mu_2 u_1 &= 2v' \sqrt{\mu_1 \mu_2 u}, \\ u_1 \sqrt{\frac{\mu_2}{\mu_1 u}} &= v \pm \sqrt{v^2 - 1} = -v' + \sqrt{v'^2 + 1}.\end{aligned}$$

v a un minimum pour $v = 1$, et la condition pour que cette valeur se trouve entre les limites de l'intégrale est

$$1 < \sqrt{\frac{\mu_1 u}{\mu_2}} < u,$$

condition qu'on peut encore exprimer en disant que μ_1 et μ_2 doivent tous les deux être plus petits que $\sqrt{\mu_1 \mu_2 u}$. u_1 étant croissant, le radical $\sqrt{v^2 - 1}$ passe, lors du minimum, du négatif au positif. Si le minimum a lieu pour une valeur moindre que la limite inférieure de l'intégrale ($\mu_1 u < \mu_2$) le signe à prendre entre les limites de l'intégrale est le signe *plus*; c'est, au contraire le signe *moins*, si la valeur qui donne lieu au minimum est plus grande que la limite supérieure de l'intégrale ($\mu_1 > \mu_2 u$).

Si l'on pose pour abréger

$$\begin{aligned}2\sqrt{\mu_1 \mu_2 u} &= a, & \mu_1 u + \mu_2 &= a v_0, & \mu_1 + \mu_2 u &= a v_1, \\ \mu_1 u - \mu_2 &= a v'_0, & \mu_1 - \mu_2 u &= a v'_1,\end{aligned}$$

on obtiendra de cette manière

$$\begin{aligned}X_2^2 &= \Sigma \Sigma \frac{4u}{a} \left[\int_{v_0}^{v_1} dv \left(1 \pm \frac{v}{\sqrt{v^2 - 1}} \right) \sin av \right. \\ &\quad \left. + \int_{v'_0}^{v'_1} dv' \left(-1 + \frac{v'}{\sqrt{v'^2 + 1}} \right) \sin av' \right].\end{aligned}$$

Ici l'on a

$$\begin{aligned} & \sum' \sum' \frac{4u}{a} \left(\int_{v_0}^{v_1} dv \sin av - \int_{v_0'}^{v_1'} dv' \sin av' \right) \\ &= 2 \sum \sum \frac{\sin \mu_1 \sin \mu_2 u - \sin \mu_1 u \sin \mu_2}{\mu_1 \mu_2}, \end{aligned}$$

expression qui s'évanouit, les deux termes devenant identiques si l'on permute les indices.

L'expression est donc réduite à

$$X_2^2 = \sum' \sum' \frac{4u}{a} \left[\int_{v_0}^{v_1} dv \left(\pm \frac{v}{\sqrt{v^2 - 1}} \right) \sin av + \int_{v_0'}^{v_1'} dv' \frac{v'}{\sqrt{v'^2 + 1}} \sin av' \right],$$

le double signe étant déterminé de la manière indiquée ci-dessus.

Nous considérerons d'abord l'intégrale

$$\int \frac{dz \cdot z}{\sqrt{z^2 - 1}} e^{azi}$$

relative à la variable complexe $z = v + wi$. Si l'intégrale est prise le long d'une courbe fermée, cette intégrale s'évanouira; si donc on intègre le long d'un rectangle dont les sommets sont situés en 1, v_0 , $v_0 + \infty i$, $1 + \infty i$, on trouvera

$$\begin{aligned} 0 = & \int_1^{v_0} \frac{dv \cdot v}{\sqrt{v^2 - 1}} e^{v a i} + i \int_0^\infty \frac{dw \cdot (v_0 + wi)}{\sqrt{(v_0 + wi)^2 - 1}} e^{a(v_0 + wi)i} \\ & - i \int_0^\infty \frac{dw (1 + wi)}{\sqrt{(1 + wi)^2 - 1}} e^{a(1 + wi)i}. \end{aligned}$$

Les deux derniers termes peuvent être développés en séries semiconvergentes; mais, pour abréger les calculs, nous ne conserverons que les termes qui ne s'évanoui-

ront pas dans le résultat final, l'expression de X_2^2 , pour $u = \infty$, et nous procéderons de la même manière dans ce qui suit. En conséquence, nous obtiendrons

$$0 = \int_1^{v_0} \frac{dv \cdot v}{\sqrt{v^2 - 1}} e^{av_i} + \frac{i v_0 e^{av_0 i}}{a \sqrt{v_0^2 - 1}} - \frac{1}{2} \sqrt{\frac{2\pi}{a}} e^{\left(a + \frac{\pi}{4}\right)i},$$

et par suite

$$\int_0^{v_0} \frac{dv \cdot v}{\sqrt{v^2 - 1}} \sin av = -\frac{v_0 \cos av_0}{a \sqrt{v_0^2 - 1}} + \frac{1}{2} \sqrt{\frac{2\pi}{a}} \sin\left(a + \frac{\pi}{4}\right).$$

Ici v_0 est supposé plus grand que 1. Dans le cas de $v_0 = 1$, il va sans dire que l'intégrale s'évanouira. Une équation analogue peut être obtenue en remplaçant v_0

* NOTE 6. par v_1 .*

Si la valeur qui donne le minimum est comprise entre les limites de l'intégrale, on aura

$$1 < \sqrt{\frac{\mu_1 u}{\mu_2}} < u,$$

$$a \sqrt{v_0^2 - 1} = \mu_1 u - \mu_2 = a v'_0,$$

$$a \sqrt{v_1^2 - 1} = \mu_2 u - \mu_1 = -a v'_1,$$

$$\int_{v_0}^{v_1} dv \left(\pm \frac{v}{\sqrt{v^2 - 1}} \right) \sin av = - \int_{v_0}^1 \frac{dv \cdot v}{\sqrt{v^2 - 1}} \sin av + \int_1^{v_1} \frac{dv \cdot v}{\sqrt{v^2 - 1}} \sin av.$$

Donc on aura, d'après l'équation trouvée ci-dessus, dans ce cas

$$\begin{aligned} & \int_{v_0}^{v_1} dv \left(\pm \frac{v}{\sqrt{v^2 - 1}} \right) \sin av \\ &= -\frac{v_0 \cos av_0}{a v'_0} + \frac{v_1 \cos av_1}{a v'_1} + \sqrt{\frac{2\pi}{a}} \sin\left(a + \frac{\pi}{4}\right), \end{aligned}$$

tandis que l'intégrale se transformera pour $v_0 = 1$ en

$$\frac{v_1 \cos av_1}{av'_1} + \frac{1}{2} \sqrt{\frac{2\pi}{a}} \sin\left(a + \frac{\pi}{4}\right),$$

et pour $v_1 = 1$ en

$$-\frac{v_0 \cos av_0}{av'_0} + \frac{1}{2} \sqrt{\frac{2\pi}{a}} \sin\left(a + \frac{\pi}{4}\right).$$

Si la valeur qui donne le minimum est moindre que la limite inférieure, on aura

$$\mu_1 u < \mu_2, \quad \sqrt{v_0^2 - 1} = -v'_0, \quad \sqrt{v_1^2 - 1} = -v'_1,$$

$$\int_{v_0}^{v_1} dv \left(\pm \frac{v}{\sqrt{v^2 - 1}} \right) \sin av = -\frac{v_0 \cos av_0}{av'_0} + \frac{v_1 \cos av_1}{av'_1},$$

et, si cette valeur est plus grande que la limite supérieure de l'intégrale, on obtiendra

$$\mu_1 > \mu_2 u, \quad \sqrt{v_0^2 - 1} = v'_0, \quad \sqrt{v_1^2 - 1} = v'_1,$$

$$\int_{v_0}^{v_1} dv \left(\pm \frac{v}{\sqrt{v^2 - 1}} \right) \sin av = -\frac{v_0 \cos av_0}{av'_0} + \frac{v_1 \cos av_1}{av'_1}.$$

En remarquant qu'on a en tous cas

$$\int_{v_0'}^{v_1'} \frac{dv \cdot v}{\sqrt{v^2 + 1}} \sin av = \frac{v'_0 \cos av'_0}{av_0} - \frac{v'_1 \cos av'_1}{av_1},$$

on verra que X_2^2 peut être exprimé par

$$X_2^2 = S_1 \frac{4u}{a} \sqrt{\frac{2\pi}{a}} \sin\left(a + \frac{\pi}{4}\right) + \frac{1}{2} S_2 \frac{4u}{a} \sqrt{\frac{2\pi}{a}} \sin\left(a + \frac{\pi}{4}\right)$$

$$+ S_3 \frac{4u}{a^2} \left[-\frac{v_0 \cos av_0}{v'_0} + \frac{v_1 \cos av_1}{v'_1} + \frac{v'_0 \cos av'_0}{v_0} - \frac{v'_1 \cos av'_1}{v_1} \right],$$

où S_1 , S_2 et S_3 désignent trois sommes doubles par rapport à m_1 et m_2 , telles que le produit $m_1 m_2$ parcourt dans S_1 et S_2 la série entière des nombres de 1 jusqu'à ∞ , tandis que chaque facteur m_1 ou m_2 ne parcourt dans S_1 que la suite des nombres plus petits que $\sqrt{m_1 m_2 u}$, et dans S_2 la suite qui correspond précisément à $m_1 = \sqrt{m_1 m_2 u}$ et $m_2 = \sqrt{m_1 m_2 u}$. Dans la somme double S_3 , m_1 et m_2 parcourent tous nombres entiers de 1 jusqu'à ∞ à l'exception des valeurs correspondantes soit à $m_1 = \sqrt{m_1 m_2 u}$,

* NOTE 7. soit à $m_2 = \sqrt{m_1 m_2 u}$.*

La dernière somme double se décompose, si l'on y porte les valeurs données de a , v_0 , v_1 , v'_0 , v'_1 , en les deux sommes

$$S_3 \frac{1}{\mu_1 \mu_2} \left(-\frac{\mu_1 u + \mu_2}{\mu_1 u - \mu_2} \cos(\mu_1 u + \mu_2) + \frac{\mu_1 u - \mu_2}{\mu_1 u + \mu_2} \cos(\mu_1 u - \mu_2) \right) \\ + S_3 \frac{1}{\mu_1 \mu_2} \left(\frac{\mu_1 + \mu_2 u}{\mu_1 - \mu_2 u} \cos(\mu_1 + \mu_2 u) - \frac{\mu_1 - \mu_2 u}{\mu_1 + \mu_2 u} \cos(\mu_1 - \mu_2 u) \right),$$

qui sont égales, comme on le reconnaît en permutant les indices. Si l'on pose ici $\mu_1 = 2\pi m_1$, $\mu_2 = 2\pi m_2$, $u = x + \frac{1}{2}$, la somme totale des deux sommes se réduira à

$$-S_3 \frac{(-1)^{m_1}}{\pi^2 m_1} \left(\frac{1}{m_1(x + \frac{1}{2}) - m_2} + \frac{1}{m_1(x + \frac{1}{2}) + m_2} \right).$$

Dans cette somme double on doit, d'après la condition de la sommation, négliger les termes qui correspondent à $m_1(x + \frac{1}{2}) - m_2 = 0$; grâce à cette condition, la sommation peut facilement être exécutée. De cette manière la somme double ci-dessus se réduit à

$$\frac{1}{\pi^2(x + \frac{1}{2})} \sum_{m_1=1}^{m_1=\infty} \frac{(-1)^{m_1}}{m_1^2} = -\frac{1}{12(x + \frac{1}{2})}.$$

En conséquence, cette somme double peut être négligée,

si l'on ne tient pas compte des termes qui s'évanouissent pour $x = \infty$. Alors le résultat total peut être exprimé par

$$X_2^2 = \frac{u}{\pi \sqrt{2}} \sum \sum \frac{\sin \left(4\pi (m_1 m_2 u)^{\frac{1}{2}} + \frac{\pi}{4} \right)}{(m_1 m_2 u)^{\frac{3}{4}}},$$

et la sommation double doit être exécutée comme il a été dit ci-dessus pour les sommations désignées par S_1 et S_2 .

Le résultat trouvé met en évidence que, si l'on néglige les termes qui s'évanouissent pour $n = \infty$, les seuls éléments utiles de l'intégrale qui représentait originairement X_2^2 seront les éléments voisins du point minimum; et, ce théorème étant acquis, le calcul peut être exécuté plus facilement d'une autre manière.

Si l'on désigne la valeur de u_1 au point minimum par v_1 , on aura $v_1 = \sqrt{\frac{\mu_1 u}{\mu_2}}$, et la condition pour que ce point soit situé entre les limites de l'intégrale sera $1 < v_1 < u$. Si cette condition est remplie, on peut poser, en dehors du point minimum $u_1 = v_1(1 + y)$, et regarder y comme assez petit pour qu'on soit en droit, dans le développement suivant les puissances de y , de négliger les puissances de y supérieures à la seconde. Enfin, l'intégrale peut être prise entre des limites indéterminées $-\omega$ et $+\omega$ (définies avec plus de précision dans mon mémoire Vid. Selsk. Skr., série 6, tome VI, p. 13, Œuvres scientifiques, tome I, p. 422) ou, ce qui ici revient au même, de $-\infty$ à $+\infty$. Si le point minimum se confond avec une limite de l'intégrale, l'une de ces limites de y sera remplacée par 0 et le résultat doit en ce cas être divisé par 2. Si le point minimum est situé en dehors des limites de l'intégrale, le résultat sera 0. De cette manière on obtiendra

$$X_s^2 = 2u\sqrt{\pi} \Sigma \Sigma \frac{\sin \left(2\sqrt{\mu_1 \mu_2 u} + \frac{\pi}{4} \right)}{(\mu_1 \mu_2 u)^{\frac{3}{4}}},$$

les sommations devant être exécutées de manière que $1 < \sqrt{\frac{\mu_1 u}{\mu_2}} < u$ et devant être divisées par 2 dans les cas où l'on a, soit $1 = \sqrt{\frac{\mu_1 u}{\mu_2}}$, soit $u = \sqrt{\frac{\mu_1 u}{\mu_2}}$.

Cette expression est précisément identique au résultat trouvé ci-dessus; et, comme nous avons de cette manière démontré la légitimité du procédé, nous en ferons une application au calcul de X_s^2 .

L'expression trouvée (17) peut être écrite sous la forme

$$X_s^2 = 2u \Sigma \Sigma' \dots \int_1^u \frac{du_1}{u_1} \dots \\ \dots \int_1^{u_{s-2}} \frac{du_{s-1}}{u_{s-1}} \frac{1}{\mu_s u_{s-1}} \sin \left(\mu_s u_{s-1} \pm \mu_{s-1} \frac{u_{s-2}}{u_{s-1}} \dots \pm \mu_1 \frac{u}{u_1} \right),$$

où les doubles signes expriment qu'on prend la somme de toutes les expressions correspondantes aux différentes combinaisons de ces signes. Cependant un minimum ou un maximum n'est possible que dans le cas où tous les signes sont positifs et il est alors déterminé par

$$\mu_s u_{s-1} = \mu_{s-1} \frac{u_{s-2}}{u_{s-1}} = \dots \mu_1 \frac{u}{u_1} = \nu,$$

où

$$\nu = (\mu_s \mu_{s-1} \dots \mu_1 u)^{\frac{1}{s}}.$$

Les conditions pour que ce point soit situé entre les limites de l'intégrale sont exprimées par

$$\mu_p < \nu \text{ pour } p = 1, 2, 3 \dots s.$$

Si les valeurs des variables $u_1, u_2, \dots, u_{s-1}$ au point minimum sont désignées par $v_1, v_2, \dots, v_{s-1}$, elles seront déterminées par

$$v_1 = \frac{\mu_1 u}{\nu}, \quad v_2 = \frac{\mu_1 \mu_2 u}{\nu^2}, \quad \dots, \quad v_{s-1} = \frac{\mu_1 \mu_2 \dots \mu_{s-1} u}{\nu^{s-1}} = \frac{\nu}{\mu_s}.$$

Nous introduirons maintenant de nouvelles variables $y_1, y_2, \dots, y_{s-1}$, définies par les équations

$$u_1 = v_1(1 + (s-1)y_1), \quad u_2 = v_2(1 + (s-2)(y_1 + y_2)), \\ \dots, \quad u_{s-1} = v_{s-1}(1 + y_1 + y_2 + \dots + y_{s-1}).$$

Ces variables peuvent être considérées comme assez petites pour être négligées dans les coefficients des fonctions trigonométriques et l'on peut arrêter le développement de l'angle

$$\mu_s u_{s-1} + \mu_{s-1} \frac{u_{s-2}}{u_{s-1}} + \dots + \mu_1 \frac{u}{u_1}$$

suivant les puissances de y_p aux termes du second degré par rapport à y_p . Dans ce développement*, les coefficients * NOTE 8. de y_p et de $y_p y_q$ s'évanouiront si p diffère de q , et le résultat se réduira à

$$\nu s + \frac{1}{2} \nu (s(s-1)y_1^2 + (s-1)(s-2)y_2^2 + \dots + 2 \cdot 1 y_{s-1}^2).$$

Ensuite, si les intégrations par rapport à $y_1, y_2, \dots, y_{s-1}$ sont étendues de $-\infty$ à $+\infty$, on obtiendra

$$X_s^s = 2u(s-1)! \Sigma \Sigma \dots \frac{1}{\nu} \int_{-\infty}^{+\infty} dy_1 \dots \\ \dots \int_{-\infty}^{+\infty} dy_{s-1} \sin\left(\nu s + \frac{1}{2} \nu (s(s-1)y_1^2 + \dots + 2 \cdot 1 y_{s-1}^2)\right) \\ = \frac{2u}{\nu^s} \Sigma \Sigma \dots \frac{1}{\nu} \left(\frac{2\pi}{\nu}\right)^{\frac{s-1}{2}} \sin\left(\nu s + (s-1)\frac{\pi}{4}\right).$$

Ici l'on a $\nu = (\mu_1 \mu_2 \dots \mu_s u)^{\frac{1}{s}} = 2\pi (m_1 m_2 \dots m_s u)^{\frac{1}{s}}$,
et par conséquent

$$X_s^s = \frac{u}{\pi \sqrt{s}} \sum \sum \dots \frac{\sin \left(2\pi s (m_1 m_2 \dots m_s u)^{\frac{1}{s}} + (s-1) \frac{\pi}{4} \right)}{(m_1 m_2 \dots m_s u)^{\frac{s+1}{2s}}}. \quad (24)$$

Dans cette sommation s -uple, le produit $m_1 m_2 \dots m_s$ varie de 1 à ∞ ; mais chaque facteur pris séparément ne doit pas, d'après la condition $\mu_p < \nu$, excéder la limite $(m_1 m_2 \dots m_s u)^{\frac{1}{s}}$ et, quand il atteint cette limite, auquel cas le point minimum est confondu avec l'une des limites de la variable u_p , le résultat ne doit être compté que pour moitié.

Si l'on pose $m_1 m_2 \dots m_s = m$, on pourra transformer la sommation s -uple en une sommation simple, pourvu qu'on multiplie l'expression sur laquelle porte la sommation par un facteur $\gamma_n^s(m)$ qui indique le nombre des décompositions différentes de m en s facteurs, y compris le nombre 1 et sous la condition qu'aucun facteur ne dépasse la limite $u = (mu)^{\frac{1}{s}}$. Dans les cas où cette limite est atteinte, les résultats ne sont comptés que pour moitié. Avec ces notations, l'équation (24) prendra la forme

$$X_s^s = \frac{u}{\pi \sqrt{s}} \sum_{m=1}^{m=\infty} \gamma_n^s(m) \frac{\sin \left(2\pi s (mu)^{\frac{1}{s}} + (s-1) \frac{\pi}{4} \right)}{(mu)^{\frac{s+1}{2s}}}, \quad (25)$$

$$n = (mu)^{\frac{1}{s}}.$$

La quantité $\gamma_n^s(m)$, définie ci-dessus, peut être déterminée comme coefficient de m^r dans le développement

$$\begin{aligned} & (1 + 2^r + 3^r \dots + (n-1)^r + \tfrac{1}{2} n^r)^s \\ &= 1 + \gamma_n^s(2) 2^r + \gamma_n^s(3) 3^r \dots + \gamma_n^s(m) m^r + \dots \end{aligned} \quad (26)$$

Si dans l'équation (10) on remplace $B^s(x)$ par la valeur de X_s^* trouvée de cette manière, en vue d'obtenir la partie périodique correspondante de $\vartheta(x)$, et si l'on effectue les calculs numériques pour une valeur de x , donnée très grande, et pour des valeurs de m choisies arbitrairement, on reconnaîtra que, pour des valeurs croissantes de s , les termes, après avoir présenté des variations de signe tout à fait irrégulières, se réunissent successivement en groupes de plus en plus grands de même signe et que c'est de la sommation de ces groupes que le résultat dépend essentiellement. On peut facilement mettre en évidence la formation d'un tel groupe.

On pose

$$2s(mu)^{\frac{1}{s}} + \frac{s-1}{4} = \theta_s;$$

alors le terme suivant sera déterminé par le développement

$$\theta_{s+1} = \theta_s + \frac{d\theta_s}{ds} \cdot \frac{1}{1} + \frac{d^2\theta_s}{ds^2} \cdot \frac{1}{1 \cdot 2} + \dots$$

Si $\frac{d^2\theta_s}{ds^2}$ et les dérivées suivantes sont des quantités très petites et si $\frac{d\theta_s}{ds}$ est un nombre entier impair, positif ou négatif, dans la série (10) où les termes ont des signes alternés, un groupe de ladite espèce se formera autour du terme correspondant à s .

On a ici

$$\begin{aligned} \frac{d\theta_s}{ds} &= 2(mu)^{\frac{1}{s}} \left(1 - \frac{\log mu}{s} \right) + \frac{1}{4}, \\ \frac{d^2\theta_s}{ds^2} &= 2(mu)^{\frac{1}{s}} \frac{(\log mu)^2}{s^2}, \end{aligned}$$

ou, si l'on pose $\frac{\log mu}{s} = y$,

$$\frac{d\theta_s}{ds} = 2e^y(1-y) + \frac{1}{4}, \quad \frac{d^2\theta_s}{ds^2} = 2e^y \frac{y^3}{\log mu}.$$

On doit donc avoir, p étant un entier positif, négatif ou nul,

$$\frac{d\theta_s}{ds} = 2e^y(1-y) + \frac{1}{4} = 1 - 2p.$$

Comme y est nécessairement positif, les valeurs négatives

* NOTE 9. de p ne sont pas admissibles, et l'on trouvera*

$$y = 0,83774\dots, \quad e^y = 2,31114\dots \text{ pour } p = 0,$$

$$y = 1,19011\dots, \quad e^y = 3,28746\dots \text{ pour } p = 1,$$

$$y = 1,40051\dots, \quad e^y = 4,05727\dots \text{ pour } p = 2,$$

$$y = 1,55459\dots, \quad e^y = 4,73317\dots \text{ pour } p = 3.$$

Puis on verra que $\frac{d^2\theta_s}{ds^2}$ et les dérivées suivantes prendront des valeurs très petites, au moins pour les valeurs les plus basses de p , si $\log mu$ est un nombre très grand. Au dessus des limites pratiques ($\log x = 21$) et pour les valeurs inférieures de m , cette condition n'est satisfaite qu'assez approximativement; pourtant, on pourra facilement faire ressortir la formation des groupes correspondants aux valeurs inférieures de p , si x est de l'ordre de grandeur des millions.

Dans ces groupes entrera donc

$$\sin \pi \theta_s = \sin \pi \left(\log mu \cdot \frac{2e^y + \frac{1}{4}}{y} - \frac{1}{4} \right),$$

et, si la fonction périodique est écrite sous la forme

$$\sin \pi \theta_s = \sin 2\pi \left(\frac{\log u}{\lambda} + k_m \right), \quad (27)$$

λ sera déterminé par

$$\lambda = \frac{y}{e^y + \frac{1}{8}}. \quad (28)$$

La partie de $\vartheta(x)$ qui correspond au groupe considéré peut donc être représentée par une courbe dont les points d'intersection successifs avec l'axe des abscisses ont une distance constante $\frac{1}{2}\lambda$, si $\log u$ ou, ce qui revient au même, $\log x$ est pris pour abscisse.

Aux valeurs indiquées ci-dessus de y correspondent $\lambda = 0,34388 \dots, 0,34875 \dots, 0,33487 \dots, 0,31999 \dots$

Toutefois on a dû remarquer que ce sont précisément les grandes valeurs de s , à savoir celles qui sont voisines de $\log mu$, qui produisent les longues périodes de $\vartheta(x)$; mais, à cause de cette circonstance, il faut de nouveau discuter le calcul qui a conduit à l'équation (25). Il faut se rappeler que nous avons étendu la variation des variables y_p à l'infini dans les deux directions, ce qui peut être légitime, à condition que y_p dépasse certaines limites étroites. Mais, si le nombre s des variables et celui des intégrales deviennent si grands qu'on ne puisse pas considérer $u^{\frac{1}{s}}$ comme un grand nombre, ce qui est précisément le cas ici, il n'est plus permis d'étendre à ce point la variation de y_p . La conséquence sera que θ_s , surtout pour les valeurs les plus petites de p , s'approche fort de sa limite inférieure, qui est $2s(mu)^{\frac{1}{s}}$ augmenté d'une constante indépendante de s . Si nous passons à cette limite inférieure elle-même, nous aurons

$$\frac{d\theta_s}{ds} = 2e^y(1-y) = 1-2p,$$

et les valeurs de y calculées par cette équation seront

$$\begin{aligned} y &= 0,76803 \dots, & e^y &= 2,15553 \dots & \text{pour } p &= 0, \\ y &= 1,15718 \dots, & e^y &= 3,18097 \dots & \text{pour } p &= 1, \\ y &= 1,37809 \dots, & e^y &= 3,96731 \dots & \text{pour } p &= 2, \\ y &= 1,53736 \dots, & e^y &= 4,65232 \dots & \text{pour } p &= 3. \end{aligned}$$

La période elle-même peut à présent être déterminée par

$$\lambda = y e^{-y} \quad (29)$$

et aux valeurs calculées de y correspondront respectivement

$$\lambda = 0,35631 \dots, \quad 0,36378 \dots, \quad 0,34736 \dots, \quad 0,33045 \dots$$

Les valeurs vraies de λ seront donc comprises entre ces valeurs et celles qu'on a calculées ci-dessus; et, pour les valeurs les plus petites de p , auxquelles correspondent les groupes les plus grands, elles seront certainement plus proches des dernières valeurs. Du reste, les périodes qui correspondent aux valeurs inférieures de p diffèrent tellement peu qu'elles se combineront sur une grande étendue en une seule période qui ne diffère que peu de 0,35.

On peut, en outre, conclure de l'expression (25) que les coefficients des fonctions périodiques trigonométriques ou les amplitudes seront, pour les valeurs les plus grandes de s , et par suite aussi de λ , à peu près du même ordre de grandeur que $x^{\frac{1}{2}}$ et d'un ordre inférieur pour les s et les λ décroissants et par suite pour les périodes courtes. La détermination plus précise de l'amplitude et de la phase des périodes différentes semble pourtant présenter des difficultés très considérables.

Le résultat auquel je suis parvenu par les calculs ci-dessus peut être retrouvé, en ce qu'il a d'essentiel, par un autre procédé, que j'exposerai ici parce qu'il me semble jeter une lumière nouvelle sur la question.

Posons

$$(1 + 2^r + 3^r + \dots)^s = 1 + \gamma^s(2) 2^r + \dots + \gamma^s(x) x^r + \dots, \quad (30)$$

$\gamma^s(x)$ désignant le nombre des décompositions différentes de x en un produit de s facteurs, y compris l'unité.

Posons de plus

$$G^s(x) = 1 + \gamma^s(2) + \dots + \gamma^s(x). \quad (31)$$

On peut calculer $\gamma^s(x)$ à l'aide de $\gamma^{s-1}(x)$ en sommant les décompositions qui correspondent au facteur nouveau; on obtient de cette manière

$$\gamma^s(x) = \sum \gamma^{s-1}\left(\frac{x}{d}\right),$$

la sommation étant étendue à tous les diviseurs de x . De plus, en se servant du symbole E de Legendre pour représenter le plus grand entier contenu dans une fraction, on a

$$\sum \gamma^{s-1}\left(\frac{x}{d}\right) = \sum_{q=1}^{q=x} \left(E\left(\frac{x}{q}\right) - E\left(\frac{x-1}{q}\right) \right) \gamma^{s-1}(q);$$

par où l'on obtiendra

$$\gamma^s(x) = \sum_{q=1}^{q=x} \left(E\left(\frac{x}{q}\right) - E\left(\frac{x-1}{q}\right) \right) \gamma^{s-1}(q), \quad (32)$$

et par suite

$$G^s(x) = \sum_{q=1}^{q=x} E\left(\frac{x}{q}\right) \cdot \gamma^{s-1}(q). \quad (33)$$

Ces expressions peuvent être mises sous forme analytique, grâce aux relations

$$E\left(\frac{x}{q}\right) - E\left(\frac{x-1}{q}\right) = \sum_{m=1}^{m=q} \frac{\cos 2\pi \frac{mx}{q}}{q}, \quad (34)$$

* NOTE 10.

$$E\left(\frac{x}{q}\right) + \frac{1}{2} = \sum_{m=1}^{m=q} \frac{\sin 2\pi \frac{m(x+\frac{1}{2})}{q}}{2q \sin \pi \frac{m}{q}}, * \quad (35)$$

dont on constate facilement l'exactitude, en effectuant les sommations indiquées.

Nous analyserons avec plus de précision $\gamma^s(x)$ sous la forme qui en résulte, à savoir

$$\left. \begin{aligned} \gamma^s(x) &= \sum_{q=1}^{q=x} \sum_{m=1}^{m=q} \frac{\cos 2\pi \frac{mx}{q}}{q} \gamma^{s-1}(q) \\ &= \sum_{m=1}^{m=x} \sum_{q=m}^{q=x} \frac{\cos 2\pi \frac{mx}{q}}{q} \gamma^{s-1}(q). \end{aligned} \right\} \quad (36)$$

Dans la dernière expression posons

$$q = q_1 q_2 \dots q_{s-1};$$

puis la sommation simple par rapport à q est transformée en une sommation $(s-1)$ -uple par rapport à $q_1, q_2, \dots, q_{s-1}$, sommation étendue à tous les cas pour lesquels

$$m \overline{\leq} q_1 q_2 \dots q_{s-1} \overline{\leq} x. \quad (37)$$

Comme cette transformation fait disparaître le facteur $\gamma^{s-1}(q)$, on obtiendra

$$\gamma^s(x) = \sum_{m=1}^{m=x} \sum \dots \sum \frac{\cos 2\pi \frac{mx}{q_1 q_2 \dots q_{s-1}}}{q_1 q_2 \dots q_{s-1}}, \quad (38)$$

les signes $\sum \sum \dots$ indiquant la sommation $(s-1)$ -uple, exécutée sous lesdites conditions par rapport aux $(s-1)$ variables $q_1, q_2, \dots, q_{s-1}$.

Nous considérerons d'abord la sommation simple

$$\sum \frac{a_1}{q_1} \cos 2\pi \frac{a_1}{q_1},$$

où q_1 parcourt une longue série de nombres successifs, série qui, du reste, n'est pas déterminée avec plus de précision, et où a_1 est un nombre très grand. On remarquera alors qu'il se formera pour certaines valeurs de q_1 que nous désignerons par q'_1 , de part et d'autre du terme correspondant à q'_1 , un groupe de termes de même signe. On trouvera qu'il en est ainsi sous les conditions suivantes.

Qu'on fasse $q_1 = q'_1 + y_1$ et qu'on suppose

$$\frac{a_1}{(q'_1 + k_1)^2} = m_1,$$

où k_1 est une fraction proprement dite, m_1 un nombre entier. On pourra alors former le développement

$$\frac{a_1}{q_1} = \frac{a_1}{q'_1 + k_1} - \frac{a_1(y_1 - k_1)}{(q'_1 + k_1)^2} + \frac{a_1(y_1 - k_1)^2}{(q'_1 + k_1)^3} - \dots,$$

qui, par une élimination partielle de k_1 et q'_1 , se transforme en

$$\frac{a_1}{q_1} = 2\sqrt{m_1 a_1} - m_1(y_1 + q'_1) + \frac{m_1^{\frac{3}{2}}(y_1 - k_1)^2}{a_1^{\frac{1}{2}}} - \dots$$

Ici $m_1(y_1 + q'_1)$ est un nombre entier; et si, de plus, m_1 est en même temps un nombre très petit en comparaison de $a_1^{\frac{1}{2}}$, les conditions pour la formation du groupe considéré seront satisfaites. Si a_1 atteint la limite $a_1 = \infty$ et si m_1 est supposé fini, les termes de chaque groupe particulier peuvent être sommés exactement en remplaçant la sommation par rapport à y_1 par une intégration entre les limites $-\infty$ et $+\infty$, pourvu que le groupe

ne soit pas un groupe limite, interrompu par une valeur limite donnée de q_1 . Le résultat de l'intégration est

$$\frac{a_1}{\sqrt{2}} \frac{\cos \left(2\pi \cdot 2(m_1 a_1)^{\frac{1}{2}} + \frac{\pi}{4} \right)}{(m_1 a_1)^{\frac{1}{2}}}.$$

La formation du groupe aura encore lieu, si a_1 est fini, mais très grand, tandis que m_1 appartient aux termes inférieurs de la série des nombres, et la somme des termes d'un groupe sera à peu près la même que celle que nous avons trouvée ci-dessus. Mais, à mesure que la grandeur de m_1 approche de $a_1^{\frac{1}{2}}$, le nombre des termes du groupe va décroître, et l'expression de la somme va perdre de son exactitude; puis la formation des groupes cessera complètement. Mais, si m_1 croît ensuite constamment et dépasse la limite $a_1^{\frac{1}{2}}$, on peut inversement considérer chaque terme à la variable q_1 comme produit par la sommation d'un groupe de termes à la variable m_1 , groupe dans lequel chaque terme a précisément la forme indiquée plus haut. On peut facilement se convaincre de l'exactitude de cette assertion par un

* NOTE 11. calcul analogue à celui qui a été exécuté ci-dessus*. Par conséquent, si l'on peut négliger les termes pour lesquels m_1 (comme q_1) s'approche de $a_1^{\frac{1}{2}}$, la transformation indiquée de la sommation relative à la variable q_1 en une sommation relative à la variable m_1 peut être considérée comme valable en général pour tous les q_1 .

Cette restriction faite par rapport à la zone neutre où aucune formation de groupe n'a lieu d'aucun côté, on pourra poser

$$\sum_{q_1} \frac{a_1}{q_1} \cos \frac{2\pi a_1}{q_1} = \sum_{m_1} \frac{a_1}{m_1} \frac{\cos \left(2\pi \cdot 2(m_1 a_1)^{\frac{1}{2}} + \frac{\pi}{4} \right)}{(m_1 a_1)^{\frac{1}{2}}}, \quad (39)$$

équation où les limites des deux sommes peuvent être déduites des relations entre les deux variables q_1 et m_1 .

Posons ensuite $\alpha_1 = \frac{a_2}{q_2}$ et concevons qu'une sommation soit effectuée dans les deux membres par rapport à la variable q_2 , supposée parcourir une longue série de nombres successifs. Dans le second membre posons $q_2 = q'_2 + y_2$ et

$$\frac{(m_1 \alpha_2)^{\frac{1}{2}}}{(q'_2 + k_2)^{\frac{3}{2}}} = m_2,$$

où k_2 est une fraction proprement dite m_2 un nombre entier. On pourra alors, comme précédemment, former le développement

$$2 \left(\frac{m_2 \alpha_2}{q_2} \right)^{\frac{1}{2}} = 3(m_1 m_2 \alpha_2)^{\frac{1}{2}} - m_2(y_2 + q'_2) + \frac{3}{4} \frac{m_2^2}{(m_1 m_2 \alpha_2)^{\frac{1}{2}}} (y_2 - k_2)^2 - \dots$$

Ici $m_2(y_2 + q'_2)$ est un nombre entier, et par conséquent un groupe se formera autour du terme correspondant à q'_2 , si m_2^2 est suffisamment petit en comparaison de $(m_1 m_2 \alpha_2)^{\frac{1}{2}}$ et par suite m_2 petit en comparaison de $(m_1 \alpha_2)^{\frac{1}{2}}$. Les termes du groupe étant sommés comme ci-dessus, on obtiendra

$$\sum \sum \frac{a_2}{q_1 q_2} \cos 2\pi \frac{a_2}{q_1 q_2} = \sum \sum \frac{a_2}{\sqrt{3}} \frac{\cos \left(2\pi \cdot 3(m_1 m_2 \alpha_2)^{\frac{1}{2}} + 2\frac{\pi}{4} \right)}{(m_1 m_2 \alpha_2)^{\frac{1}{2}}},$$

équation dont la validité peut être étendue comme celle de l'équation (39) à des valeurs de m_2 très grandes, tandis qu'elle cesse de subsister quand m_2 et q_2 approchent de $(m_1 \alpha_2)^{\frac{1}{2}}$.

De cette manière on parviendra sans difficulté à une formule valable en général avec les mêmes restrictions que ci-dessus, savoir

$$= \left. \begin{aligned} & \sum \sum \dots \frac{1}{q_1 q_2 \dots q_p} \cos 2\pi \frac{a_p}{q_1 q_2 \dots q_p} \\ & \sum \sum \dots \frac{1}{V^{p+1}} \frac{\cos \left(2\pi (p+1) (m_1 m_2 \dots m_p a_p)^{\frac{1}{p+1}} + p \frac{\pi}{4} \right)}{(m_1 m_2 \dots m_p a_p)^{\frac{p}{2p+2}}} \end{aligned} \right\} \quad (40)$$

Les relations entre les deux systèmes de variables peuvent approximativement être exprimées par

$$\frac{a_1}{q_1^{\frac{1}{2}}} = m_1, \quad \frac{(m_1 a_2)^{\frac{1}{2}}}{q_2^{\frac{3}{2}}} = m_2, \quad \frac{(m_1 m_2 a_3)^{\frac{1}{3}}}{q_3^{\frac{4}{3}}} = m_3, \quad \dots$$

où

$$a_1 = \frac{a_2}{q_2}, \quad a_2 = \frac{a_3}{q_3}, \quad \dots$$

On peut déduire de là les équations approchées

$$m_1 q_1 = m_2 q_2 = \dots = m_p q_p = (m_1 m_2 \dots m_p a_p)^{\frac{1}{p+1}}. \quad (41)$$

Si l'on en fait application aux sommations indiquées dans la formule (38), on doit poser $p = s-1$ et $a_{s-1} = mx$; mais on doit toutefois remarquer qu'on ne peut, de cette manière, représenter qu'imparfaitement l'expression totale de $\gamma^s(x)$, tant à cause des zones où la formation de groupes cesse en même temps des deux côtés, qu'à cause de la détermination des limites des variables. Cependant, ce qui importe ici, c'est de tenir compte séparément des termes de $\gamma^s(x)$ dont dépend la formation des longues périodes de la série des nombres premiers: et, pour cet objet, c'est surtout la détermination des limites inférieures des nouvelles variables $m_1, m_2, \dots$, qui a de l'importance. Il résulte de l'équation (41), quand on y fait $p = s-1$, et $a_{s-1} = mx$, qu'on a approximativement

$$q_1 q_2 \dots q_{s-1} = \frac{m x}{(m m_1 \dots m_{s-1} x)^{\frac{1}{s}}}, \quad (42)$$

et d'après (37) la limite supérieure de ce produit est x . Toutes les variables $m_1, m_2 \dots m_{s-1}$ ont donc 1 pour limite inférieure, tant que m est petit et ne dépasse point la limite $x^{\frac{1}{s-1}}$.

Nous pouvons donc, dans l'expression (38) de $\gamma^s(x)$, choisir une suite de termes qui peuvent être transformés en une somme s -uple

$$\sum \sum \dots \frac{1}{V^s} \frac{\cos \left(2\pi \cdot s (m m_1 \dots m_{s-1} x)^{\frac{1}{s}} + (s-1) \frac{\pi}{4} \right)}{(m m_1 \dots m_{s-1} x)^{\frac{s-1}{2s}}}, \quad (43)$$

où toutes les variables parcourent la série des nombres, depuis 1 jusqu'à certaines limites que nous ne déterminons pas.

Si nous cherchons ensuite à déterminer la série correspondante de $G^s(x)$, nous poserons d'abord

$$G^s(x) = \sum_{x'=x_0}^{x'=x} \gamma^s(x') + G^s(x_0 - 1),$$

où x_0 et x peuvent être considérés l'un et l'autre comme des nombres très grands; et, si nous remplaçons ici $\gamma^s(x')$ par l'expression (43), nous pourrions effectuer, au lieu de la sommation par rapport à x' , une intégration dont le résultat est facile à obtenir approximativement, en remarquant que toutes les valeurs que parcourt x' sont de grands nombres. Prise séparément, cette partie de $G^s(x)$ qui dépend de x , peut être exprimée par *

* NOTE 12.

$$\frac{1}{2} \cdot \frac{x}{\pi V^s} \sum \sum \dots \frac{\sin \left(2\pi \cdot s (m m_1 \dots m_{s-1} x)^{\frac{1}{s}} + (s-1) \frac{\pi}{4} \right)}{(m m_1 \dots m_{s-1} x)^{\frac{s+1}{2s}}}. \quad (44)$$

La forme de cette expression est identique, sauf le facteur $\frac{1}{2}$, à celle de X_s^s que nous avons trouvée dans l'équation (24). Dans les deux cas, 1 est la limite inférieure de toutes les s variables; mais la ressemblance ne va pas plus loin, les deux développements représentant deux fonctions périodiques différentes, quoique d'une grande affinité.

Comme on a, eu égard à la série (9),

$$A^s(x) = G^s(x) - \frac{s}{1} G^{s-1}(x) + \frac{s(s-1)}{1 \cdot 2} G^{s-2}(x) - \dots + (-1)^s,$$

on peut facilement obtenir un développement de $\vartheta(x)$ analogue à (10), de la forme

$$\vartheta(x) = -b_0 + b_1 \frac{G^1(x)}{1} - b_2 \frac{G^2(x)}{2} + \dots \pm b_{s_1} \frac{G^{s_1}(x)}{s_1},$$

$$s_1 > \frac{\log x}{\log 2} - 1.$$

Si l'on y porte l'expression de $G^s(x)$ donnée par (44), on obtiendra la partie correspondante de $\vartheta(x)$. Comme on le verra facilement, la discussion ultérieure de cette expression sera essentiellement la même que précédemment; le résultat sera pareillement de constater l'existence de grandes périodes ayant $\log x$ comme variable et la période λ définie par l'équation (28). Mais, en ce qui concerne la détermination plus précise des amplitudes et des phases des fonctions périodiques, les mêmes obstacles se présenteront encore ici, quoique sous une forme nouvelle.

J'ai déjà mentionné dans l'introduction que les écarts périodiques entre les nombres de nombres premiers

trouvés effectivement et ceux qu'on a calculés par la formule de Riemann semblent se développer successivement pour les grands nombres en périodes de plus en plus régulières.

Heureusement, il se trouve que M. le Dr. Gram a traité cette question, d'un point de vue empirique, dans son mémoire cité dans l'introduction. Il s'explique sur cette question de la manière suivante (p. 250):

„Glaisher, pour mieux mettre en évidence la variation des écarts, en a fait une représentation graphique dans un diagramme qu'il a joint à son mémoire. Il y a quelque chose qui pourrait faire soupçonner une période dépendante de $\ln$, en ce sens que la période est à peu près 0,17, $\log_{10} n$ étant pris pour argument, et par conséquent 0,39, $\ln$ étant l'argument.“

Par une communication verbale M. Gram m'a donné le moyen de compléter cette indication. La fonction périodique étant mise sous la forme

$$\sin 2\pi \left(\frac{\log x}{\lambda} + C \right),$$

M. Gram a déterminé les constantes par une compensation géométrique simple et il a trouvé

$$\lambda = 0,39, \quad C = 0,34.$$

Cette formule a bien manifesté des écarts considérables avec le diagramme du neuvième million; mais M. Gram a négligé ces écarts parce qu'il a soupçonné qu'ils provenaient en partie d'erreurs que contiendraient les tableaux des nombres premiers calculés par Dase*; en * NOTE 13. effet, les résultats de ce dénombrement sont difficiles à mettre d'accord avec le calcul du nombre des nombres

premiers jusqu'à 10 millions fait par Meissel. Au contraire, la formule a mis en évidence une grande concordance avec le diagramme de 3 jusqu'à 8 millions.

Enfin, après avoir calculé les écarts avec la formule de Riemann pour la série des nombres premiers correspondante à l'intervalle de $\log x = 0,1$ jusqu'à $\log x = 15$ ($x = 3269017$) M. Gram est parvenu au résultat suivant: „Que l'apparente distribution régulière des grands maxima et minima du diagramme de Glaisher est vraisemblablement due à un accident.“

Tant que la tentative de déterminer les périodes ne comportait que des tâtonnements, nécessités par les données relativement peu nombreuses qu'on possédait et dont une partie, en outre, était entachée du soupçon d'inexactitude, il pouvait sembler sage de renoncer à cette tentative. Heureusement, M. Gram a fait connaître les résultats de ses essais de compensation; et l'on reconnaît à présent, à la lumière de la théorie, que la double conjecture d'après laquelle on devait prendre $\log x$ pour l'argument et 0,39 pour la valeur de la période, était essentiellement correcte et qu'une discussion ultérieure des matériaux empiriques dont nous disposons jusqu'ici est tout à fait superflue. La circonstance même que les longues périodes régulières s'effacent, si x ne dépasse pas une certaine limite inférieure, devient une confirmation de la théorie.

Le fondement théorique obtenu est un fort encouragement à poursuivre la détermination exacte du nombre des nombres premiers. Il serait surtout du plus haut intérêt de contrôler le calcul, fait par Dase, du nombre des nombres premiers jusqu'à 9 millions, soit qu'on se

servît de la méthode de Meissel, soit qu'on fit usage d'une méthode analogue. Il importerait aussi de déterminer une suite de points entre 10 et 100 millions, par exemple de 10 en 10 millions. On fournirait de cette manière des matériaux précieux pour la continuation des recherches théoriques.

NOTES.

NOTE 1. On reconnaît facilement qu'on a

$$\beta^s(x) = \alpha^s(x) + \frac{s}{2} \alpha^{s-1}(x) + \frac{s(s-1)}{2 \cdot 4} \alpha^{s-2}(x) + \dots$$

ou, sous forme symbolique, $\beta^s(x) \simeq (\alpha + \frac{1}{2})^s(x)$.

On en déduit

$$\alpha^s(x) \simeq (\beta - \frac{1}{2})^s(x),$$

et par conséquent

$$A^s(x) \simeq (B - \frac{1}{2})^s(x).$$

NOTE 2. Dans cette note je chercherai à reconstruire le raisonnement de Lorenz sur lequel est fondée la formule (14).

Soit donnée la série

$$S = \int_{x_1}^{x_2} f(x) \sum (1 + 2 \sum_1^{\infty} \cos 2\pi m_1 x) dx,$$

où $f(x)$ est une fonction finie et continue entre les limites x_1 et x_2 ; quelle est la valeur de cette série?

Posons $s_m = 1 + 2 \sum_1^m \cos 2\pi m x$; alors on aura, x n'étant pas un entier

$$s_m = \frac{\sin(2m+1)q}{\sin q}, \quad q = \pi x,$$

$$S = \lim_{m=\infty} \int_{x_1}^{x_2} f(x) \frac{\sin(2m+1)q}{\sin q} dx,$$

et, si $\frac{f(x)}{\sin q} = \phi(\chi)$, $\phi(x)$ sera encore finie et continue entre les limites x_1 et x_2 si aucun entier n'est compris entre ces limites.

En intégrant par parties, on obtiendra

$$\int_{x_1}^{x_2} \frac{f(x) \sin(2m+1)q}{\sin q} dx$$

$$= - \left| \frac{f(x) \cos(2m+1)q}{(2m+1)\pi \sin q} \right|_{x_1}^{x_2} + \frac{1}{(2m+1)\pi} \int_{x_1}^{x_2} \frac{d\phi(x)}{dx} \cos(2m+1)q dx.$$

Par conséquent on aura, si m croît à l'infini,

$$\lim \int_{x_1}^{x_2} \frac{f(x) \sin(2m+1)q}{\sin q} dx = 0 = S.$$

Si les entiers

$$r, \quad r+1, \quad r+2 \dots (r+s)$$

sont compris entre les limites x_1 et x_2 , on peut diviser l'intégrale de la manière suivante

$$S = \left(\int_{x_1}^{r-\delta} + \int_{r-\delta}^{r+\delta'} + \int_{r+\delta'}^{r+1-\delta_1} + \int_{r+1-\delta_1}^{r+1+\delta'_1} + \dots + \int_{r+s-\delta_s}^{r+s+\delta'_s} \int_{r+s+\delta'_s}^{x_2} \right) \varphi(x) dx,$$

$$\varphi(x) = f(x) (1 + 2 \sum_1^{\infty} \cos 2m_1 \pi x),$$

les quantités δ étant arbitraires, mais aussi petites qu'on veut.

D'après ce qui précède, les intégrales $\int_{x_1}^{r-\delta}$, $\int_{r+\delta'}^{r+1-\delta_1}$,
 $\int_{r+1+\delta'_1}^{r+2-\delta_2}$... s'évanouiront. Reste à déterminer les intégrales

$\int_{r+p-\delta_p}^{r+p+\delta_p} \varphi(x) dx$, p étant un nombre entier compris entre les nombres r et $r+s$. On aura, en remplaçant $r+p$ par r' ,

$$\begin{aligned} \int_{r'-\delta_p}^{r'+\delta_p} \varphi(x) dx &= \int_{r'-\delta_p}^{r'+\delta_p} f(x) (1 + 2 \sum_1^{\infty} \cos 2\pi m_1 x) dx \\ &= \int_{-\delta_p}^{+\delta_p} f(r' + x) (1 + 2 \sum_1^{\infty} \cos 2\pi m_1 x) dx \\ &= \int_0^{\delta_p} f(r' + x) (1 + 2 \sum_1^{\infty} \cos 2\pi m_1 x) dx \\ &\quad + \int_0^{\delta_p} f(r' - x) (1 + 2 \sum_1^{\infty} \cos 2\pi m_1 x) dx. \end{aligned}$$

Par des méthodes bien connues on reconnaîtra que les deux intégrales sont égales à $\frac{f(r)}{2}$.

Par conséquent

$$S = f(r) + f(r+1) \dots f(r+s),$$

à l'exception des cas où l'une des limites de l'intégrale ou toutes les deux sont égales à des entiers. Dans ces cas les termes correspondants doivent être comptés pour moitié.

Considérons maintenant la formule (14) de Lorenz. Dans ce qui précède, Lorenz dit qu'il cherchera une expression de la puissance s -ième de $\frac{1}{2} + 2^r + 3^r \dots + x^r$. Mais il ne le fait pas et il n'a pas besoin d'une telle expression. Au contraire il trouve par la formule (14) une expression intégrale de

$$B^s(x) = \beta^s(1) + \beta^s(2) \dots \beta^s(x),$$

et la remarque d'après laquelle tous les éléments de l'intégrale s -uple qui se trouve au second membre de l'équation (14) s'évanouiront, à moins que

$$\frac{u_1}{u_2}, \frac{u_2}{u_3} \dots \frac{u_{s-1}}{u_s}, \frac{u_s}{1}$$

ne soient des entiers, met cette proposition en évidence. C'est pourquoi nous chercherons avec plus de détails à vérifier cette remarque.

Soit donnée une intégrale double

$$\int_1^u du_1 \int_1^{u_1} f(u_1, u_2) du_2.$$

Comme on le reconnaît facilement, on peut intervertir l'ordre des intégrations. L'intégrale est une somme double de tous les éléments $du_1 du_2$ multipliés par $f(u_1, u_2)$ sous les conditions que u_1 et u_2 soient tous deux plus grands que 1, que u_1 soit plus grand que u_2 et que u_1 soit plus petit que u . C'est pourquoi l'on aura

$$\int_1^u du_1 \int_1^{u_1} f(u_1, u_2) du_2 = \int_1^u du_2 \int_{u_2}^u f(u_1, u_2) du_1.$$

On reconnaît de la même manière que

$$\begin{aligned} & \int_1^u du_1 \int_1^{u_1} du_2 \dots \int_1^{u_{s-1}} f(u_1, u_2 \dots u_s) du_s \\ &= \int_1^u du_s \int_{u_s}^u du_{s-1} \int_{u_{s-1}}^u du_{s-2} \dots \int_{u_2}^u f(u_1, u_2 \dots u_s) du_1, \end{aligned}$$

et par suite

$$\begin{aligned}
& \int_1^u du_1 \int_1^{u_2} \frac{du_2}{u_2} \dots \int_1^{u_{s-1}} \frac{du_s}{u_s} \left(1 + 2 \sum' \cos \mu_1 \frac{u_1}{u_2} \right) \dots \left(1 + 2 \sum' \cos \mu_{s-1} \frac{u_{s-1}}{u_s} \right) \\
& \quad (1 + 2 \sum' \cos \mu_s u_s) \\
& = \int_1^u \frac{du_s}{u_s} (1 + 2 \sum' \cos \mu_s u_s) \int_{u_s}^u \frac{du_{s-1}}{u_{s-1}} \left(1 + 2 \sum' \cos \mu_{s-1} \frac{u_{s-1}}{u_s} \right) \\
& \quad \dots \int_{u_3}^u \frac{du_2}{u_2} \left(1 + 2 \sum' \cos \frac{\mu_2 u_2}{u_3} \right) \int_{u_2}^u du_1 \left(1 + 2 \sum' \cos \frac{\mu_1 u_1}{u_2} \right).
\end{aligned}$$

En conséquence de la proposition développée ci-dessus, tous les éléments de la dernière intégrale s'évanouiront, à moins que $\frac{u_1}{u_2}$ ne soit un nombre entier. Dans ce cas l'élément sera égal à u_2 (ou à $\frac{u_2}{2}$, si c'est un élément limite). En conséquence on aura

$$\begin{aligned}
& \int_{u_3}^u \frac{du_2}{u_2} \left(1 + 2 \sum' \cos \frac{\mu_2 u_2}{u_3} \right) \int_{u_2}^u du_1 \left(1 + 2 \sum' \cos \frac{\mu_1 u_1}{u_2} \right) \\
& = \int_{u_3}^u du_2 \left(1 + 2 \sum' \cos \frac{\mu_2 u_2}{u_3} \right)
\end{aligned}$$

sous la condition que $\frac{u_1}{u_2}$ soit un entier. En continuant de cette manière, on verra que chaque élément de l'intégrale s -uple donnée est égal à zéro, à moins que toutes les fractions

$$\frac{u_1}{u_2}, \quad \frac{u_2}{u_3} \dots \frac{u_{s-1}}{u_s}, \quad \frac{u_s}{1}$$

ne soient des entiers. Dans ces cas, les éléments sont égaux à l'unité ou, si ce sont des éléments limites, à une puissance de $\frac{1}{2}$.

Comme on le reconnaît facilement, l'exposant de cette puissance de $\frac{1}{2}$ indique le nombre des fractions $\frac{u_1}{u_2}, \frac{u_2}{u_3} \dots$ qui sont égales à l'unité. Ces remarques mettent en évidence l'exactitude de la formule (14).

NOTE 3. La démonstration de l'équation (15) met en évidence qu'on peut permuter les trois fonctions f, g, h . Dès lors on peut transporter la somme $2 \sum \cos \mu_s u_s$, comme toutes les autres sommes, de droite à gauche.

NOTE 4. Les constantes C_n peuvent être définies par l'équation

$$C_n = \frac{(-1)^n}{n!} \lim_{p=\infty} \int_1^p \frac{du}{u} (\log u)^n 2 \sum \cos \mu u$$

où p est un entier.

On aura alors ($n > 0$),

$$C_n = \lim_{p=\infty} \frac{(-1)^n}{n!} \left(\frac{\log^n 2}{2} + \frac{\log^n 3}{3} + \dots + \frac{\log^n (p-1)}{p-1} + \frac{\log^n p}{2p} - \frac{\log^{n+1} p}{n+1} \right).$$

Comme on le voit facilement, ces valeurs sont, au signe près, les coefficients de la série

$$s \zeta(1-s) = c_0 + c_1 s + c_2 s^2 + \dots$$

où ζ (la fonction de Riemann) est définie par l'équation

$$\zeta(s) = \lim_{p=\infty} \left(\sum_1^p p^{-s} - \frac{p^{1-s}}{1-s} \right).$$

Ces coefficients ont été calculés par M.M. J.-L.-V.-W. Jensen et Gram. Voir Comptes rendus 1887, tome 104 p. 1157 et Overs. over det kgl. danske Vidensk. Selsk. Forhandling 1895, p. 308.

NOTE 5. Si l'on pose

$$z = u\phi(z) + v,$$

où ϕ est une fonction donnée, u et v des variables indépendantes, on aura, $f(z)$ étant une fonction arbitraire,

$$f(z) = \lambda_0 + \frac{\lambda_1 u}{1} + \frac{\lambda_2 u^2}{1 \cdot 2} + \frac{\lambda_3 u^3}{1 \cdot 2 \cdot 3} + \dots$$

$$\lambda_0 = f(v), \quad \lambda_n = \frac{d^{n-1}(\phi(v)^n f'(v))}{dv^{n-1}}.$$

Pour faire des applications de cette formule, on doit encore chercher le reste; mais, dans ce qui suit, nous supposons que ce reste tend vers zéro pour un nombre croissant de termes.

Nous posons

$$z = -uz + v,$$

$$f(z) = \int_a^z \frac{e^{-2z}}{z} dz, \quad f'(z) = \frac{e^{-2z}}{z}.$$

La série citée de Lagangre nous donnera alors

$$\begin{aligned} \int_a^z \frac{e^{-2z}}{z} dz &= \int_a^v \frac{e^{-2z}}{z} dz - \frac{uv e^{-2v}}{v} + \frac{u^2}{1 \cdot 2} \frac{d\left(\frac{v^2 e^{-2v}}{v}\right)}{dv} \\ &\quad - u^3 \frac{d^2\left(\frac{v^3 e^{-2v}}{v}\right)}{dv^2} + \dots \end{aligned}$$

Différentions cette équation par rapport à v ; nous obtiendrons ainsi

$$\begin{aligned} \frac{e^{-2z}}{z} \frac{dz}{dv} &= \frac{e^{-2v}}{v} - \frac{2v}{1+v} \frac{e^{-2v}}{v} - u \frac{d\left(\frac{v^2 e^{-2v}}{v}\right)}{dv} + \frac{u^2}{1 \cdot 2} \frac{d^2\left(\frac{v^3 e^{-2v}}{v}\right)}{dv^2} \\ &\quad - u^3 \frac{d^3\left(\frac{v^4 e^{-2v}}{v}\right)}{dv^3} + \dots \end{aligned}$$

Si nous posons $u = 2C_0$, nous aurons la formule (23) de Lorenz.

NOTE 6. Considérons l'intégrale

$$\int_0^{\infty} \frac{dw (v_0 + wi)}{V(v_0 + wi)^2 - 1} e^{a(v_0 + wi)i}$$

où a est très grand et où v_0 diffère de l'unité. Alors ce ne sont que les petites valeurs de w qui influent sur la valeur de l'intégrale et nous pouvons sans erreur sensible remplacer l'intégrale donnée par l'intégrale plus simple

$$\frac{v_0 e^{av_0 i}}{Vv_0^2 - 1} \int_0^{\infty} dw e^{-aw} = \frac{1}{a} \frac{v_0 e^{av_0 i}}{Vv_0^2 - 1}.$$

De la même manière, on peut remplacer

$$\int_0^{\infty} \frac{dw (1 + wi)}{V(1 + wi)^2 - 1} e^{a(1 + wi)i}$$

par

$$e^{ai} \int_0^{\infty} \frac{dw e^{-aw}}{V2wi} = -ie^{(a + \frac{\pi}{4})i} V \frac{\pi}{2a}.$$

De cette façon on obtient

$$0 = \int_1^{v_0} \frac{dv}{Vv^2 - 1} + \frac{iv_0 e^{av_0 i}}{aVv_0^2 - 1} - \frac{1}{2} V \frac{2\pi}{a} e^{(a + \frac{\pi}{4})i};$$

mais cette équation n'est pas valable quand v_0 est égal à 1 ou ne diffère que peu de 1. Dans ce cas le membre $\frac{v_0 e^{av_0 i}}{aVv_0^2 - 1}$ de l'équation devient infini, tandis que l'intégrale

$\int_1^{v_0} \frac{dv}{Vv^2-1}$ est égale à zéro; c'est pourquoi l'on ne sait pas si l'expression trouvée peut être appliquée à la sommation approchée de la série qui exprime X_s^2 .

NOTE 7. L'indication de Lorenz qu'on doit dans la somme double S_s négliger les termes qui correspondent à $m_1 = \sqrt{m_1 m_2} u$ et $m_2 = \sqrt{m_1 m_2} u$ n'est pas correcte: on ne doit négliger que les termes pour lesquels les dénominateurs s'évanouissent. D'ailleurs Lorenz lui-même n'a dans les sommations suivantes négligé que ces termes.

NOTE 8. Posons

$$S_p = y_1 + y_2 \dots y_p;$$

alors, en négligeant tous les termes d'un ordre supérieur au second, nous aurons

$$\begin{aligned} U &= \mu_s u_{s-1} + \mu_{s-1} \frac{u_{s-2}}{u_{s-1}} + \dots + \mu_1 \frac{u}{u_1} \\ &= \nu \left(1 + S_{s-1} + \frac{1+2S_{s-2}}{1+S_{s-1}} + \frac{1+3S_{s-3}}{1+2S_{s-2}} + \dots + \frac{1+(s-1)S_1}{1+(s-2)S_2} + \frac{1}{1+(s-1)S_1} \right) \\ &= \nu \left(\begin{array}{l} s-1 \cdot 2 S_{s-2} S_{s-1} - 2 \cdot 3 S_{s-3} S_{s-2} \dots - (s-2)(s-1) S_1 S_2 \\ + 1^2 \cdot S_{s-1}^2 \quad + 2^2 \cdot S_{s-2}^2 \quad \quad + (s-1)^2 S_1^2 \end{array} \right). \end{aligned}$$

Si l'on porte dans la seconde série $S_p^2 = (y_p + S_{p-1}) S_p$, $p = s-1, s-2 \dots$, on aura

$$U = \nu (s + S_{s-1} (y_{s-1} - S_{s-2}) + 2 S_{s-2} (2 y_{s-2} - S_{s-3}) \dots (s-1)^2 y_1^2).$$

Ici l'on pose $S_{s-1} = y_{s-1} + S_{s-2}$, d'où

$$U = \nu (s + y_{s-1}^2 + S_{s-2} (4 y_{s-2} - S_{s-2} - 2 S_{s-3}) + \dots).$$

Dans cette expression on pose $S_{s-2} = y_{s-2} + S_{s-3}$, d'où

$$U = \nu (s + y_{s-1}^2 + 3 y_{s-2}^2 + 3 S_{s-3} (3 y_{s-3} - S_{s-3} - S_{s-4}) + \dots).$$

En continuant de cette manière on obtiendra l'expression de Lorenz.

NOTE 9. M. J.-P. Gram m'a communiqué qu'il existe une certaine relation entre les valeurs de y trouvées ici et les racines de $\xi(t) = 0$, ξ étant la fonction de Riemann. Voici la communication de M. Gram.

L'équation

$$2\pi e^y(1-y) + \frac{1}{4} = 1 - 2p$$

met en évidence que $2\pi e^y$ doit approximativement être égal à l'une des racines, α , de $\xi(t) = 0$, puisque le nombre des racines égales à N ou plus petites que N est égal à

$$\frac{N}{2\pi} \left(\log \frac{N}{2\pi} - 1 \right) + \frac{5}{4} + \varepsilon$$

(cfr. la note de M. Gram, Vidensk. Selsk. Overs. 1902, p. 13. Note sur les zéros de la fonction $\xi(s)$ de Riemann).

Les valeurs calculées de y (p. 556) donneront

$2\pi e^y = 14,521$	tandis que	$\alpha = 14,135$
20,655		21,022
25,492		25,011
29,739		30,425.

Les valeurs de y (p. 557) donneront

$$\begin{aligned} 2\pi e^y &= 13,543 \\ &19,987 \\ &24,927 \\ &29,231. \end{aligned}$$

C'est vraisemblablement inexact quand Lorenz suppose qu'il ait trouvé les limites de λ . Il est vraisem-

blable que son y doit être précisément égal à $\log \frac{\alpha}{2\pi}$. S'il en est ainsi, le développement de Lorenz met en évidence le fait intéressant qu'il traite les racines α de $\xi(t) = 0$ sans recourir aux fonctions ξ en indiquant leur importance — à son point de vue, bien entendu, — pour la partie périodique de $\vartheta(x)$.

NOTE 10. L'équation (35) n'est pas juste, à moins que $\frac{\sin 2\pi \frac{m(x + \frac{1}{2})}{q}}{\sin \pi \frac{m}{q}}$ ne soit égal à $2(x + \frac{1}{2})$ pour $m = q$.

C'est ce qu'on reconnaît de la manière suivante.

Si l'on pose

$$S = \sum_{m=1}^{m=q-1} \frac{\sin \frac{2\pi m(x + \frac{1}{2})}{q}}{2q \sin \frac{\pi m}{q}},$$

on aura

$$S = \frac{1}{q} \sum_{m=1}^{m=q-1} \left(\frac{1}{2} + \cos \frac{2\pi m}{q} + \cos \frac{2 \cdot 2\pi m}{q} \dots \cos \frac{2x\pi m}{q} \right),$$

et

$$\sum_{m=1}^{m=q-1} \cos \frac{2 \cdot r \pi m}{q} = q - 1$$

si q est un facteur de r , tandis qu'on aura dans le cas contraire

$$\sum_{m=1}^{m=q-1} \cos \frac{2r\pi m}{q} = -1.$$

Par conséquent

$$S = \frac{1}{q} \left(\frac{q-1}{2} - x + q E\left(\frac{x}{q}\right) \right) = \frac{1}{2} + E\left(\frac{x}{q}\right) - \frac{1}{q} \left(x + \frac{1}{2} \right).$$

NOTE II. On pose $m_1 = m'_1 + z$, où

$$m'_1 = \frac{a_1}{(q_1 + k_1)^2} = \frac{a_1}{q_1^2} - \frac{2 a_1 k_1}{q_1^3} + \dots$$

Si k_1 est très petit en comparaison de q_1 , on peut dans la série m'_1 négliger les termes suivants et poser

$$\begin{aligned} \sqrt{m_1 a_1} &= \sqrt{\frac{a_1^2}{q_1^2} - \frac{2 a_1^2 k_1}{q_1^3} + z a_1} \\ &= \frac{a_1}{q_1} + \frac{1}{2} q_1 \left(z - \frac{2 a_1 k_1}{q_1^3} \right) - \frac{1}{8} \frac{q_1^3}{a_1} \left(z - \frac{2 a_1 k_1}{q_1^3} \right)^2 \\ &= \frac{1}{2} \frac{a_1}{q_1} + \frac{1}{2} q_1 (z + m'_1) - \frac{1}{8} \frac{q_1^3}{a_1} \left(z - \frac{2 a_1 k_1}{q_1^3} \right)^2. \end{aligned}$$

On aura donc avec une certaine approximation, $z + m'_1$ étant un nombre entier,

$$\begin{aligned} &\frac{a_1}{\sqrt{2}} \frac{\cos \left(2\pi \cdot 2(m_1 a_1)^{\frac{1}{2}} + \frac{\pi}{4} \right)}{(m_1 a_1)^{\frac{1}{4}}} \\ &= \frac{a_1}{\sqrt{2}} \frac{\cos 2\pi \left(\frac{a_1}{q_1} - \frac{1}{4} \frac{q_1^3}{a_1} \left(z - \frac{2 a_1 k_1}{q_1^3} \right)^2 + \frac{1}{8} \right)}{\sqrt{\frac{a_1}{q_1}}}. \end{aligned}$$

Si l'on prend la somme d'un grand nombre de pareilles expressions et si l'on remplace la somme par une intégrale (de $-\infty$ à $+\infty$), on obtiendra

$$\begin{aligned} &\Sigma \frac{a_1 \cos \left(2\pi \cdot 2(m_1 a_1)^{\frac{1}{2}} + \frac{\pi}{4} \right)}{\sqrt{2} (m_1 a_1)^{\frac{1}{4}}} \\ &= \frac{\sqrt{q_1 a_1}}{\sqrt{2}} \int_{-\infty}^{+\infty} \cos 2\pi \left(\frac{a_1}{q_1} - \frac{1}{4} \frac{q_1^3}{a_1} z^2 + \frac{1}{8} \right) dz = \frac{a_1}{q_1} \cos \frac{2\pi a_1}{q_1}. \end{aligned}$$

Je suppose que c'est de cette manière que Lorenz a prouvé son hypothèse.

Comme on le voit, le procédé ne met pourtant pas en évidence le degré d'exactitude du résultat obtenu.

NOTE 12. Lorenz dit qu'on peut approximativement poser

$$G^s(x) = \Sigma \Sigma' \dots \frac{1}{V^s} \int_{x_0}^x \frac{\cos \left(2\pi s (m m_1 \dots m_{s-1} x)^{\frac{1}{s}} + (s-1) \frac{\pi}{4} \right)}{(m m_1 \dots m_{s-1} x)^{\frac{s-1}{2s}}} dx \\ + G^s(x_0 - 1).$$

Mais on aura

$$\frac{1}{V^s} \int_{x_0}^x \frac{\cos \left(2\pi s (m m_1 \dots m_{s-1} x)^{\frac{1}{s}} + (s-1) \frac{\pi}{4} \right)}{(m m_1 \dots m_{s-1} x)^{\frac{s-1}{2s}}} dx \\ = \frac{1}{2\pi V^s} \frac{x \sin \left(2\pi s (m m_1 \dots m_{s-1} x)^{\frac{1}{s}} + (s-1) \frac{\pi}{4} \right)}{(m m_1 \dots m_{s-1} x)^{\frac{s+1}{2s}}} + C \\ - \frac{(s-1)}{2\pi \cdot 2s V^s} \int_{x_0}^x \frac{\sin \left(2\pi s (m m_1 \dots m_{s-1} x)^{\frac{1}{s}} + (s-1) \frac{\pi}{4} \right)}{(m m_1 \dots m_{s-1} x)^{\frac{s+1}{2s}}} dx$$

où C est indépendant de x .

Dans cette expression, le dénominateur de la dernière intégrale est très grand en comparaison de celui de l'intégrale proposée; aussi cette intégrale peut-elle être négligée. De cette manière on obtient l'expression (44).

NOTE 13. Voir la note de M. Gram: Rapport sur quelques calculs entrepris par M. Bertelsen et concernant les nombres premiers. *Acta mathem.*, tome 17, p. 301—314; 1893.

La revision faite par M. Bertelsen des tableaux de Dase confirme complètement la supposition de M. Gram.

Date Due

[illegible]

Carnegie Mellon University Libraries



3 8482 01069 8831

510.8 L860

v.2

~~500.756~~
Lorenz

Oeuvres scientifiques

~~500.756~~
v.2

510.8 L860

Carnegie Institute of Technology
Library
Pittsburgh, Pa.

UNIVERSAL
LIBRARY



130 083

UNIVERSAL
LIBRARY